



Joint super-resolution/segmentation approaches for the tomographic images analysis of the bone micro-structure

Yufei Li

► To cite this version:

Yufei Li. Joint super-resolution/segmentation approaches for the tomographic images analysis of the bone micro-structure. Signal and Image processing. Université de Lyon, 2018. English. NNT : 2018LYSEI125 . tel-02115267

HAL Id: tel-02115267

<https://theses.hal.science/tel-02115267>

Submitted on 30 Apr 2019

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



N° d'ordre NNT : 2018LYSEI125

THÈSE DE DOCTORAT DE L'UNIVERSITÉ DE LYON
opérée au sein de
l'INSA LYON

École Doctorale ED160
ELECTRONIQUE, ELECTROTECHNIQUE, AUTOMATIQUE

Spécialité de doctorat :
Discipline : Traitement du signal et de l'image

Soutenue publiquement/à huis clos le 20/12/2018, par :
Yufei LI

**Joint super-resolution/segmentation
approaches for the tomographic images
analysis of the bone micro-structure**

Devant le jury composé de :

KOUAME Denis, Professeur des Universités, IRIT,
SOUSSEN Charles, Professeur des Universités, Centrale-Supelec
CHAPPARD Christine, CR1, Inserm, B2OA
TALBOT Hugues, professeur associé, ESIEE

Rapporteur
Rapporteur
Examinatrice
Examineur

SIXOU Bruno, Maître de conférence, HDR, INSA-LYON
PEYRIN Françoise, DR Inserm, INSA-LYON, ESRF

Directeur de thèse
Co-directrice de thèse

Département FEDORA – INSA Lyon - Ecoles Doctorales – Quinquennal 2016-2020

SIGLE	ECOLE DOCTORALE	NOM ET COORDONNEES DU RESPONSABLE
CHIMIE	<u>CHIMIE DE LYON</u> http://www.edchimie-lyon.fr Sec. : Renée EL MELHEM Bât. Blaise PASCAL, 3e étage secretariat@edchimie-lyon.fr INSA : R. GOURDON	M. Stéphane DANIELE Institut de recherches sur la catalyse et l'environnement de Lyon IRCELYON-UMR 5256 Équipe CDFA 2 Avenue Albert EINSTEIN 69 626 Villeurbanne CEDEX directeur@edchimie-lyon.fr
E.E.A.	<u>ÉLECTRONIQUE,</u> <u>ÉLECTROTECHNIQUE,</u> <u>AUTOMATIQUE</u> http://edeea.ec-lyon.fr Sec. : M.C. HAVGOUDOUKIAN ecole-doctorale.eea@ec-lyon.fr	M. Gérard SCORLETTI École Centrale de Lyon 36 Avenue Guy DE COLLONGUE 69 134 Écully Tél : 04.72.18.60.97 Fax 04.78.43.37.17 gerard.scorletti@ec-lyon.fr
E2M2	<u>ÉVOLUTION, ÉCOSYSTÈME,</u> <u>MICROBIOLOGIE, MODÉLISATION</u> http://e2m2.universite-lyon.fr Sec. : Sylvie ROBERJOT Bât. Atrium, UCB Lyon 1 Tél : 04.72.44.83.62 INSA : H. CHARLES secretariat.e2m2@univ-lyon1.fr	M. Philippe NORMAND UMR 5557 Lab. d'Ecologie Microbienne Université Claude Bernard Lyon 1 Bâtiment Mendel 43, boulevard du 11 Novembre 1918 69 622 Villeurbanne CEDEX philippe.normand@univ-lyon1.fr
EDISS	<u>INTERDISCIPLINAIRE</u> <u>SCIENCES-SANTÉ</u> http://www.ediss-lyon.fr Sec. : Sylvie ROBERJOT Bât. Atrium, UCB Lyon 1 Tél : 04.72.44.83.62 INSA : M. LAGARDE secretariat.ediss@univ-lyon1.fr	Mme Emmanuelle CANET-SOULAS INSERM U1060, CarMeN lab, Univ. Lyon 1 Bâtiment IMBL 11 Avenue Jean CAPELLE INSA de Lyon 69 621 Villeurbanne Tél : 04.72.68.49.09 Fax : 04.72.68.49.16 emmanuelle.canet@univ-lyon1.fr
INFOMATHS	<u>INFORMATIQUE ET</u> <u>MATHÉMATIQUES</u> http://edinfomaths.universite-lyon.fr Sec. : Renée EL MELHEM Bât. Blaise PASCAL, 3e étage Tél : 04.72.43.80.46 Fax : 04.72.43.16.87 infomaths@univ-lyon1.fr	M. Luca ZAMBONI Bât. Braconnier 43 Boulevard du 11 novembre 1918 69 622 Villeurbanne CEDEX Tél : 04.26.23.45.52 zamboni@maths.univ-lyon1.fr
Matériaux	<u>MATÉRIAUX DE LYON</u> http://ed34.universite-lyon.fr Sec. : Marion COMBE Tél : 04.72.43.71.70 Fax : 04.72.43.87.12 Bât. Direction ed.materiaux@insa-lyon.fr	M. Jean-Yves BUFFIÈRE INSA de Lyon MATEIS - Bât. Saint-Exupéry 7 Avenue Jean CAPELLE 69 621 Villeurbanne CEDEX Tél : 04.72.43.71.70 Fax : 04.72.43.85.28 jean-yves.buffiere@insa-lyon.fr
MEGA	<u>MÉCANIQUE, ÉNERGÉTIQUE,</u> <u>GÉNIE CIVIL, ACOUSTIQUE</u> http://edmega.universite-lyon.fr Sec. : Marion COMBE Tél : 04.72.43.71.70 Fax : 04.72.43.87.12 Bât. Direction mega@insa-lyon.fr	M. Jocelyn BONJOUR INSA de Lyon Laboratoire CETHIL Bâtiment Sadi-Carnot 9, rue de la Physique 69 621 Villeurbanne CEDEX jocelyn.bonjour@insa-lyon.fr
ScSo	<u>ScSo*</u> http://ed483.univ-lyon2.fr Sec. : Viviane POLSINELLI Brigitte DUBOIS INSA : J.Y. TOUSSAINT Tél : 04.78.69.72.76 viviane.polsinelli@univ-lyon2.fr	M. Christian MONTES Université Lyon 2 86 Rue Pasteur 69 365 Lyon CEDEX 07 christian.montes@univ-lyon2.fr

Acknowledgement

First of all, I would like to present my deep gratitude toward my families for their continuous supports and encouragement.

Then, I'm profoundly indebted to my mentors Bruno SIXOU and Françoise PEYRIN. Their patient guide seems like a lighthouse, helps me to locate my position and where to go. They provide me a free research environment and encourage me to do what I believe. At the same time, they steer me in the right direction when they think it's necessary. When I encounter difficulties, they are always trying to help me to find solutions. Their rigorous, diligence, sense of responsibility, altruism and so many shining characteristics have and will have a strong impact on me in my life.

Afterwards, I would like to acknowledge the China Scholarship Council for their finance of this thesis. Their grant enables me to concentrate on research work. I thank Bei Jing Jiao Tong University for recommending me as a candidate for this thesis. Moreover, this work was performed within the framework of the LABEX PRIMES (ANR-11-LABX-0063) of Université de Lyon, within the program "Investissements d'Avenir" (ANR-11-IDEX-0007) operated by the French National Research Agency (ANR). Thank NVDDIA for their donation of GPU card for our academic research.

Meanwhile, thanks a lot to Andrew Burghardt for offering us experimental images in this study. Thank Alina TOMA for her contributions in image registration and delicate work in former research, Pierre-Jean Gouttenoire for his help for morphological analysis, Matthieu Martin for his assist in coding and research, Pierre Ferrier, Fabrice BELLET, Philippe Delachartre, Sebastien Vallette, François Varray, for their supports with material devices. I would like to say that the discussions with Boliang YU, Juan ABASCAL, Denis FRIBOULET, Denis BUJOREANU, Claire MOUTTON are very delightful and give me a lot of inspirations.

Last but not the least, I'm grateful to our laboratory CREATIS and all the members in the lab. This is an agreeable place for research.

Merci!

Résumé étendu

L'ostéoporose est une maladie caractérisée par une perte de masse osseuse et la dégradation de la micro-architecture osseuse. L'ostéoporose est l'une des maladies des os les plus courantes parmi les gens âgés. Même si l'ostéoporose n'est pas une maladie fatale, des fractures qu'elle occasionne peuvent induire de graves complications (lésions des vaisseaux et des nerfs, infections, raideurs), parfois accompagnées de risques de mortalité.

La densité minérale des os (BMD) est un critère pour diagnostiquer l'ostéoporose. Mais la BMD seule n'est pas suffisante pour quantifier la qualité de l'os. La micro-architecture joue un rôle important pour le diagnostic de l'ostéoporose. Elle permet d'extraire certains paramètres qui peuvent qualifier la structure osseuse. La caractérisation précise de la microstructure osseuse est ainsi cruciale pour son diagnostic.

La tomographie périphérique à haute résolution (HR-pQCT) est actuellement l'une des meilleures techniques de tomodensitométrie pour l'étude de la micro-architecture osseuse en 3D. Elle permet de scanner la micro-architecture osseuse *in vivo* sur les sites périphériques humains, typiquement avec une taille de voxel de $82\mu\text{m}$. La résolution spatiale de cette technique varie dans la gamme $130\mu\text{m}$ - $150\mu\text{m}$. Cependant, une telle résolution spatiale est encore insuffisante parce que la taille des travées est entre 100 - $200\mu\text{m}$. Par conséquent, il est difficile d'avoir une estimation précise des paramètres osseux qui sont calculés après segmentation de l'os à partir du fond. La micro-CT (μ -CT), qui peut fournir des images à une résolution spatiale beaucoup plus élevée, se limite à un examen *ex vivo*.

Dans cette thèse, nous tentons de résoudre un problème de super résolution/segmentation joint, afin d'améliorer la qualité des images HR-pQCT. Les images HR-pQCT seront considérées comme des images à basse résolution et les images μ -CT du même spécimen avec une taille de voxel de $24\mu\text{m}$, seront considérées comme des images à haute résolution et utilisées comme une vérité-terrain pour la micro-architecture osseuse. Il faut noter que la taille des voxels des images de haute résolution est de $41\mu\text{m}$ jusqu'au chapitre 6, ce qui se fait par recalage. L'objectif du problème de super résolution/segmentation jointe est d'obtenir une bonne caractérisation structurelle de l'os avec les images de super résolution.

Ce manuscrit est organisé comme suit :

- chapitre 1 : Contexte et motivation.
- chapitre 2 : Problème de super résolution et fondation mathématique.
- chapitre 3 : Variation Totale et détermination du noyau de convolution.
- chapitre 4 : Optimisation nonconvexe et nonlisse pour renforcer le contraste de l'image avec un double puit de potentiel.
- chapitre 5 : Super résolution avec un apprentissage de dictionnaire semi-couplé.
- chapitre 6 : Revue de l'apprentissage profond pour la super résolution.

chapitre 7 : Super-résolution/segmentation jointe avec une méthode de deep learning, le facteur de suréchantillonnage est de 2 et 3.42.

Le premier chapitre est consacré à la présentation du contexte médical : nous détaillons l'importance de la micro-architecture trabéculaire et les différentes techniques de tomodensitométrie, puis, nous donnons les définitions des paramètres osseux pour quantifier la qualité de l'os trabéculaire. Ensuite, nous discutons des avantages et des inconvénients des images HR-pQCT et montrons les motivations de cette thèse.

Le deuxième chapitre résume l'état de l'art des méthodologies de super-résolution et quelques notions mathématiques sur les méthodes d'optimisation. Les méthodes pour résoudre un problème de super résolution peuvent être classifiées en 4 branches: (1) interpolation dans l'espace spatial, (2) traitement des images dans le domaine fréquentiel, (3) optimisation des fonctions avec des termes de régularisation, (4) machine learning. Dans cette thèse, nous nous concentrons sur les deux dernières méthodologies. Après l'introduction de l'état de l'art, nous présentons quelques notions mathématiques à propos de l'optimisation convexe et des algorithmes optimisant des fonctions convexes ou nonconvexes qui seront utilisés dans les chapitres suivants.

Le troisième chapitre passe en revue les travaux antérieurs basés sur la régularisation par variation totale (TV) et décrit les ensembles de données d'images expérimentales disponibles. Ensuite, puisque le noyau de convolution est inconnu, nous présentons une de nos contributions pour la détermination du noyau de convolution HR-pQCT sur la base de l'information mutuelle. Déterminer le noyau avec l'information mutuelle est une idée originale, cependant, nos résultats expérimentaux montrent que le noyau estimé a des effets similaires à celui obtenu avec la minimisation de la norme L_2 .

Le quatrième chapitre traite de l'application des méthodes d'optimisation non convexes et non lisses dans notre problème de super résolution/segmentation. Puisque les histogrammes de l'image μ -CT de référence présentent une distribution bimodale, nous avons considéré une approche variationnelle combinant la régularisation TV avec un potentiel de double puits pour notre problème de super résolution. Nous avons ensuite mis en place trois schémas numériques différents pour résoudre ce problème d'optimisation non convexe et non lisse, l'un d'entre eux étant basé sur la généralisation de l'algorithme 'méthode de direction alternée de multiplicateur'. La conclusion est que le potentiel de double puits améliore le contraste de l'image super résolution, mais ne parvient pas à récupérer les détails des images.

Le cinquième chapitre examine l'application des méthodes d'apprentissage par dictionnaire dans le problème de la super résolution/segmentation. Notre méthode est basée sur l'approche semi-coupled dictionary learning dans laquelle deux dictionnaires sont appris correspondant à la haute et la basse résolution. Certaines améliorations sont proposées pour traiter les grandes images osseuses 3D. Les résultats numériques montrent que l'image obtenue surpasse les méthodes précédentes. Certains critères (Dice ou BV/TV) sont sensiblement améliorés, mais la restauration de la densité de connectivité de la structure osseuse reste un grand défi.

Le sixième chapitre passe en revue les techniques de super résolution à apprentissage profond et présente l'application de la super résolution à apprentissage profond dans le traitement des images médicales. L'architecture évolue d'une simple structure à un réseau dense et profond. Le réseau d'apprentissage profond a une capacité étonnante

d'approximation de fonctions complexes. L'adaptation du réseau d'apprentissage en profondeur pour le traitement d'images médicales est encore un domaine de recherche actif.

Le chapitre sept présente les résultats préliminaires sur l'apprentissage en profondeur des méthodes de super résolution pour améliorer encore la qualité de l'image. Dans la méthode du dictionnaire, les coefficients des représentations des images haute et basse résolution sur les deux dictionnaires sont reliés linéairement par l'intermédiaire de deux matrices. Les méthodes d'apprentissage profond permettent de découvrir des fonctions plus complexes. Nous avons appliqué différentes variantes de méthodes accélérées de réseaux de neurones convolutionnels à la super résolution avec le GPU Nvidia. Dans ce chapitre, nous résolvons ce problème inverse avec un facteur de suréchantillonnage égal à 2 et 3.42. Les résultats sont très prometteurs.

En conclusion, nous avons abordé un problème de super résolution/segmentation joint appliqué aux images médicales expérimentales de tomodensitométrie osseuse. Nous avons d'abord envisagé des méthodes variationnelles pour estimer le noyau inconnu du problème et une approche non convexe et non lisse pour améliorer le contraste des résultats obtenus avec la régularisation TV. Ensuite, nous avons considéré les approches de l'apprentissage dictionnaire ainsi que les approches deep learning, qui surpassent les méthodes précédentes et nos analyses montrent qu'ils sont donc prometteurs pour l'avenir.

Chapitre 1 : Contexte et motivations

L'ostéoporose est l'une des maladies osseuses les plus courantes auxquelles les gens doivent faire face aujourd'hui. Elle augmente la fragilité de l'os et augmente le risque de fracture. Le taux de fracture augmente avec l'âge en raison de la dégradation de la micro-architecture osseuse et de la perte de la masse osseuse. Les gens ont tendance à perdre de la masse osseuse après 50 ans, surtout les femmes après la ménopause. L'ostéoporose n'est pas une maladie mortelle, mais les fractures qu'elle induit peuvent avoir de graves complications (lésions des vaisseaux et des nerfs, infections, raideur), parfois même des risques mortels. Une fois la fracture survenue, les patients ont besoin d'une longue période de soins infirmiers pour se rétablir, ce qui représente également un lourd fardeau en termes de qualité de vie personnelle et d'économie.

Les trois principaux sites de fracture chez les humains sont la hanche, les vertèbres et l'avant-bras. La fracture de la hanche nécessite 30 jours d'hospitalisation en moyenne pendant lesquels le taux de morbidité et de mortalité augmente. Entre-temps, seulement 30% des patients retrouvent leur qualité de vie d'origine. Et les patients qui ont souffert de fractures de la hanche ont un risque plus élevé de souffrir d'une autre fracture de la hanche plus tard (Silman, 1995). Les fractures vertébrales nécessitent rarement une hospitalisation. Elles ont un impact plus long sur la qualité de vie du patient plutôt au lieu d'augmenter la mortalité comme la fracture de la hanche. La fracture de l'avant-bras est moins nocive que les fractures de la hanche ou des vertèbres. Cependant, elle peut aussi endommager les nerfs et les vaisseaux environnants ou provoquer d'autres complications, et il faut de 3 à 6 mois pour guérir complètement. Les patients après une fracture de l'avant-bras peuvent avoir un risque plus élevé d'une autre fracture plus tard.

La densité minérale osseuse (DMO) est un critère de diagnostic de l'ostéoporose et est mesurée par absorptiométrie aux rayons X à double énergie (DXA). Le DXA dispose de deux faisceaux de rayons X d'énergie différente. La DMO peut être estimée en éliminant l'absorption des tissus mous. Cependant, cette technique ne peut fournir une estimation *in vivo* précise puisqu'elle ne peut mesurer des objets qu'avec deux substances avec des différences d'absorption des rayons X, ce qui n'est pas le cas pour l'imagerie *in vivo*. Les recherches ont montré que la DMO ne représente que 70 à 75% de la variabilité de la résistance des os, le reste étant lié à la micro-architecture osseuse et à la composition des tissus.

Conventionnellement, l'os est classé en deux catégories : l'os compact (cortical) et l'os trabéculaire (spongieux). Les os corticaux sont très compacts et denses. Ils forment la couche externe très compacte de l'os. Les os trabéculaires forment la couche interne de l'os avec une densité plus faible. Les os trabéculaires se régénèrent généralement plus rapidement que les os corticaux, et sont donc plus sensibles à la perte de masse osseuse. L'imagerie de la micro-architecture trabéculaire permet d'estimer de nombreux paramètres osseux cruciaux, y compris la masse osseuse, de sorte que la quantification de la qualité osseuse devient plus précise. De tels paramètres caractérisant la force de la micro-architecture osseuse facilitent et améliorent le diagnostic et le traitement de l'ostéoporose ou d'autres maladies du squelette.

HR-pQCT est un dispositif conçu pour les images 3D *in vivo* de la micro-architecture osseuse sur les sites périphériques de l'homme. Il permet aux chercheurs d'analyser quantitativement l'os *in vivo*, favorise le développement du diagnostic clinique ainsi que le

suivi du traitement. La résolution des scanners HR-pQCT est encore limitée, une résolution plus élevée est attendue à l'avenir. La tomographie est un système d'imagerie non destructif pour les échantillons osseux et non invasif. Malgré l'impossibilité de scanner les tissus humains *in vivo*, la μ -CT est une excellente technique pour scanner *in vivo* de petits animaux et des échantillons de cadavres *ex vivo* ainsi que des biopsies. Puisque la μ -CT peut fournir des paramètres précis dans l'analyse topologique et morphologique, il est généralement considéré comme un étalon pour l'amélioration des mesures HR-pQCT. HR- μ -CT a une résolution encore meilleure que la μ -CT et a de nombreuses applications pour l'étude de micro-structures de l'ordre du μm ou à l'échelle nanométrique. Dans cette thèse, nous traitons des images obtenues par HR-pQCT et μ -CT.

Des paramètres osseux sont proposés pour caractériser la micro-architecture osseuse. BV/TV reflète le rapport entre le volume osseux et le volume total. La densité de connectivité est un indicateur topologique normalisé par le volume total décrivant la structure osseuse. La structure trabéculaire est composée d'éléments en forme de plaques et de tiges. La modification de la structure osseuse est associée à la variation du nombre d'éléments en forme de plaques et de tiges. L'indice du modèle de structure, l'épaisseur trabéculaire, la séparation trabéculaire et d'autres paramètres ont été proposés pour décrire les caractéristiques morphologiques de l'os.

Même si les techniques μ -CT ou HR- μ -CT fournissent des images CT avec une très bonne résolution, elles sont limitées à l'imagerie *ex vivo*. HR-pQCT donne accès à l'imagerie *in vivo* au niveau des sites périphériques distaux : tibias et radius.

Dans cette thèse, nous tentons d'améliorer la résolution de l'image HR-pQCT *in vivo* afin que les caractérisations structurales se rapprochent de celles obtenues avec la technique μ -CT. Le problème initial est de reconstruire des images super résolues avec une taille de voxel de $41\mu\text{m}$. Dans le dernier chapitre, nous proposons une approche qui produit des images super résolues de taille de voxel $24\mu\text{m}$. Notre base de données disponible est composée de paires d'images haute et basse résolution qui sont scannées à partir de μ -CT et de la première génération HR-pQCT sur les mêmes échantillons.

Chapitre 2 : Super résolution et Fondation mathématique

Dans ce chapitre, nous visons à présenter l'état de l'art des problèmes de super résolution, puis nous introduisons certaines notions mathématiques qui seront utilisées dans les chapitres suivants.

Problème inverse

Soit $\mathcal{A}$ un opérateur linéaire ou non linéaire entre deux espaces $\mathcal{X}$ et $\mathcal{Y}$, $\mathcal{A} : \mathcal{X} \rightarrow \mathcal{Y}$. Pour un $\mathbf{f}^* \in \mathcal{X}$ donné, nous observons une sortie $\mathbf{g} \in \mathcal{Y}$, le problème direct est formulé comme suit:

$$\mathbf{g} = \mathcal{A}\mathbf{f}^* \quad (1)$$

Les espaces peuvent être des espaces de dimension finie, des espaces de Hilbert aux dimensions infinies comme les espaces L_2 ou les espaces de Banach. Dans les systèmes d'observation réels, les images bruitées $\mathbf{g}$ sont ce que nous observons, et nous essayons habituellement de déterminer $\mathbf{f} \in \mathcal{X}$ de sorte que $\mathbf{f}$ se rapproche le plus possible de la solution vraie $\mathbf{f}^*$. Ces problèmes sont appelés problèmes inverses. Un problème inverse bien posé satisfait à trois conditions :

1. La solution existe
2. La solution est unique
3. L'opérateur $\mathcal{A}^{-1}$ est continu

Si au moins l'une des trois conditions n'est pas satisfaite, le problème devient un problème inverse mal posé. La super résolution d'image est un problème inverse mal-posé. Si nous désignons par $\mathbf{n}$ le bruit additionnel, $\mathbf{g}$ les données observées, $\mathcal{A}$ l'opérateur de dégradation incluant les effets de sous-échantillonnage et de floutage, le problème super résolution est formulé ainsi:

$$\mathbf{g} = \mathcal{A}\mathbf{f} + \mathbf{n} \quad (2)$$

Notre objectif est de récupérer une image résolue $\mathbf{f}$ à partir d'une image basse résolution $\mathbf{g}$ de sorte que $\mathbf{f}$ soit proche de $\mathbf{f}^*$.

L'état de l'art de la super résolution

La détermination d'une image de super résolution avec plusieurs images connues en basse résolution concerne la super résolution multiframe. Le super résolution d'une image unique est un problème dans lequel uniquement une image de basse résolution est disponible. Pendant cette thèse, nous nous intéressons à la résolution de ce dernier problème : le super résolution d'une image simple.

L'interpolation spatiale estime le point interpolé en tenant compte de ses voisins, tels que les voisins les plus proches, en utilisant des polynômes ou des fonctions plus complexes comme les splines. Outre l'interpolation spatiale, la transformation d'une image dans un domaine différent est une autre approche pour résoudre le problème de super résolution. La transformation pourrait être basée sur la transformée cosinus discrète, la transformée

de Fourier, la transformée en ondelettes discrètes. L'idée de principe est d'améliorer la résolution de l'image en manipulant les coefficients obtenus avec la transformation.

Une autre classe de méthodes pour résoudre les problèmes de super résolution est basée sur la régularisation. La solution est obtenue en minimisant une fonction de régularisation incluant un terme des données et un terme de régularisation, correspondant à des informations a priori sur les propriétés de la solution, comme la sparsité ou la régularité. Dans la littérature, la fonction de régularisation est une fonctionnelle écrite comme suit :

$$\arg \min_f ||\mathcal{A}f - g||_2^2 + \mu \mathcal{R}(f) \quad (3)$$

μ est un paramètre de régularisation, équilibrant les effets du terme d'attache aux données et du terme de régularisation. La norme L_2 est la distance la plus populaire pour le terme d'attache aux données dans la littérature, alors que d'autres distances ont été étudiées aussi bien, comme l'information mutuelle ou la divergence de Kullback-Leibler. Le terme de régularisation est crucial pour la qualité de la solution : il intègre l'information préalable dans le résultat final et peut réduire les risques d'overfitting.

La régularisation de Tikhonov est basée sur la norme L_2 . Il s'agit d'un terme de régularisation convexe et différentiable qui a été largement utilisé pendant plusieurs décennies, mais son inconvénient est qu'il peut trop lisser la solution. La régularisation Lasso est basée sur la norme L_1 . Lasso est un terme de régularisation convexe mais non différentiable. Tikhonov donne de bons résultats pour une sélection de groupe, sa solution est moins sparse que celle obtenue avec Lasso. Lasso fournit des solutions sparses, donc n'est pas approprié pour une sélection de groupe.

La régularisation par variation totale (TV) lisse les images en améliorant la sparsité de la variation de l'image. La variation de l'image est définie comme la norme L_1 du gradient. C'est un terme de régularisation non différentiable. Il a une performance adéquate pour le débruitage de l'image, en particulier pour les images constantes par morceaux. La régularisation TV calcule les dérivés du premier ordre des images, des dérivées d'ordre supérieur ont également été étudiées.

L'apprentissage par dictionnaire permet de reconstruire des images par la représentation sparse sur une base d'atomes du dictionnaire. La sparsité de la représentation peut être favorisée par des approches de régularisation. Dans l'apprentissage par dictionnaire, l'ensemble des vecteurs de base est nommé dictionnaire, et chaque vecteur de base est nommé atome. Les atomes peuvent être définis par un ensemble de fonctions de base connues, telles que des filtres orientables, ou être appris à partir d'échantillons d'apprentissage.

Il est nécessaire d'apprendre deux dictionnaires pour résoudre un problème de super résolution avec l'approche d'apprentissage par dictionnaire. Un dictionnaire est appris à partir d'un ensemble de données d'images à basse résolution, l'autre est appris à partir d'images à haute résolution. Les deux dictionnaires peuvent être appris de différentes façons, mais ils doivent être reliés entre eux. Une fois les deux dictionnaires disponibles, il est possible de reconstruire des images de super résolution à partir des images basse résolution.

Récemment, les approches deep learning ont connu un grand succès dans le traitement de l'image avec des tâches telles que la classification, la segmentation, le débruitage. En fait, l'apprentissage par dictionnaire et le deep learning font partie de l'apprentissage

machine. Le réseau du deep learning est un réseau de neurones multicouches. Sa première couche est nommée comme couche d'entrée, la dernière couche est nommée comme couche de sortie. Les couches intermédiaires sont nommées couches cachées. Chaque neurone contient une opération linéaire et une fonction d'activation non linéaire. Le réseau de neurones convolutifs (CNN) est un réseau classique de deep learning. Chaque neurone représente une opération de convolution et une opération non linéaire. L'opération de convolution permet au réseau de détecter la même caractéristique dans différentes régions de l'image. Il semble que les filtres d'une couche CNN détectent les caractéristiques locales de l'image, tandis que les CNN multicouche permettent d'analyser globalement l'image et de résumer les visions locales précédentes. De plus, le CNN réduit le nombre de paramètres en les partageant entre les neurones du réseau, ce qui entraîne une réduction considérable de la mémoire.

L'application de CNN dans le problème de super résolution comprend deux étapes : concevoir le réseau et apprendre les paramètres à l'intérieur du réseau, reconstruire l'image de super résolution avec le réseau appris. La structure du réseau définit le nombre de couches, le nombre de filtres dans chaque couche et le type de fonction d'activation. Les paramètres du réseau sont les filtres de l'opération de convolution. Ce sont ces filtres qui extraient les caractères des images. Lors de l'apprentissage, les images à basse résolution sont transmises aux couches d'entrée. Avec les filtres initialisés au hasard, le réseau produit une image de super résolution estimée. La distance entre l'estimation et les images haute résolution correspondantes sera optimisée par rapport aux filtres du réseau. Le processus d'apprentissage sera arrêté lorsque le réseau convergera. Par la suite, l'image de super résolution peut être estimée en introduisant les images de basse résolution dans le réseau.

Fondation mathématique

Dans un problème d'optimisation convexe, si la fonction objectif est quadratique et l'ensemble de ses contraintes sont linéaires, il est possible de trouver sa solution analytique. Lorsque le problème d'optimisation comprend des contraintes d'équations linéaires et une fonction objectif différentiable, on peut appliquer la méthode de Newton pour trouver sa solution par itération. Toutefois, les fonctions peuvent ne pas être différentiables et convexes. Dans le chapitre 2, nous présentons l'opérateur proximal et la méthode de direction alternée de multiplicateur (ADMM) qui sont des outils très pratiques pour l'optimisation convexe, ensuite nous présentons certaines approches utilisées pour résoudre des problèmes non convexes.

L'opérateur proximal peut être considéré comme une méthode d'Euler avec un schéma implicite. ADMM se base au départ sur un Lagrangien qui est limité aux fonctions différentiables. Ensuite, la fonction Lagrangienne augmentée va au delà en introduisant un terme quadratique. Cependant, ce terme quadratique fait que les variables ne peuvent pas être mises à jour en parallèle. Par conséquent, ADMM a été proposé pour mettre à jour les variables séquentiellement.

L'optimisation convexe a un rôle important dans les problèmes d'optimisation non convexes. Elle peut fournir une initialisation pour l'optimisation globale, ou servir de limite inférieure à la solution de l'optimisation globale. Nous présentons brièvement quelques approches d'optimisation non convexes qui résolvent des problèmes non convexes : approches proximales ; transformation de fonctions non convexes en fonction convexe en

ajoutant un terme quadratique pour que la fonction devienne convexe ; obtention d'une fonction convexe par transformation conjuguée de Fenchel.

Les méthodes proximales pour les problèmes non convexes optimisent la partie nonconvexe par une descente de gradient et minimisent la partie non différentiable par l'opérateur proximal. Différents moyens ont été proposés pour accélérer l'algorithme. Certains algorithmes ajoutent un terme quadratique pour que les fonctionnelles nonconvexes deviennent convexes et modifient ce terme au cours des itérations. Les méthodes basées sur les différences de fonctions convexes résolvent le problème d'optimisation non convexe via la transformation conjuguée de Fenchel en utilisant le fait que le biconjugué de la fonction non convexe s'avère être son enveloppe convexe.

Chapitre 3 : Variation totale et la détermination des noyaux

Dans un système d'observation d'images, l'image vraie est généralement brouillée par un ou plusieurs noyaux convolutifs. Comment récupérer des images dégradées par des noyaux convolutifs est ce que l'on nomme un problème de déconvolution.

L'état de l'art de la déconvolution

Dans les problèmes de déconvolution, les origines du flou sont diverses et diverses méthodes ont été développées pour estimer les noyaux de flou (Reeves and Mersereau, 1992, Fergus et al., 2006, Joshi et al., 2008). J.Reeves *et al.* a utilisé une méthode de validation croisée pour déterminer les paramètres du flou (Reeves and Mersereau, 1992). Fergus *et al.* a tenté d'utiliser l'approche bayésienne pour retrouver les images individuelles dégradées par le tremblement de la caméra (Fergus et al., 2006). Le principe est de trouver le noyau via le gradient de l'image observée et la vérité terrain en supposant que le noyau de dégradation est une combinaison linéaire de distributions exponentielles. Joshi *et al.* a proposé de prédire d'abord les images latentes en affinant les bords des images floues, puis de déterminer le noyau de dégradation avec ces images et l'image d'entrée (Joshi et al., 2008). Les images finales de déconvolution sont obtenues avec les noyaux déterminés via l'algorithme Lucy-Richardson. De plus, la détermination du flou peut être régularisée avec une approximation de dimension finie (Denis et al., 2015) ou des méthodes basées sur des patches (Sun et al., 2013). Certaines méthodes spécifiques pour les images floues ont été étudiées dans (Cannon, 1976, Chang et al., 1991).

Méthode

Dans ce chapitre, nous proposons d'utiliser un terme des données qui est basé sur l'information mutuelle pour estimer le noyau et l'optimiser avec une approche de descente en gradient. Les noyaux obtenus avec cette approche et l'optimisation de la norme L_2 sont ensuite comparés pour résoudre un problème de super résolution sur des données réelles avec la régularisation TV. Leur efficacité pour améliorer la qualité d'images *in vivo* HR-pQCT est évaluée en comparant des paramètres quantitatifs de la micro-structure osseuse.

Expériences numériques

Les tests pour la détermination du noyau de flou ont été effectués sur des images expérimentales 3D d'os trabéculaire obtenues à partir de scanners HR-pQCT (Boutroy et al., 2005) et μ -CT (Burghardt et al., 2011). Les tailles de voxel des images HR-pQCT et μ -CT sont respectivement $82\mu\text{m}$ et $41\mu\text{m}$. Les images à haute et basse résolution testées ont été respectivement recalées à $200 \times 200 \times 200$ voxels et $100 \times 100 \times 100$ voxels. Le noyau initial $\mathcal{K}_0$ est choisi comme un noyau Gaussien avec un écart-type de 4. Plusieurs tailles de support pour le noyau ont été étudiées avec des nombres impairs allant de 5 à 13. Différents paramètres de régularisation ont également été testés. Dans la super résolution basée sur une régularisation TV, les paramètres sont choisis pour obtenir la

meilleure diminution de la fonctionnelle de régularisation, $\mu = 10$, $\beta = 7$ pour les deux noyaux. Les itérations sont arrêtées lorsque

$$\frac{\|\mathbf{f}^{n+1} - \mathbf{f}^n\|_2}{\|\mathbf{f}^{n+1}\|_2} \leq 10^{-3} \quad (4)$$

Résultats

Les images de super-résolution obtenues ont des structures plus minces que les images basse résolution mais les structures reconstruites ont toujours une masse osseuse plus importante que l'image haute résolution. Les deux méthodes améliorent la qualité de l'image pour le rapport BV/TV, mais ce ratio est encore trop élevé par rapport aux images à haute résolution. La densité de connectivité est un critère topologique. Même si elle a été améliorée, la divergence finale de sa courbe d'évolution montre que les noyaux déterminés n'ont pas récupéré de manière appropriée les structures topologiques dégradées. En général, le noyau déduit de la norme L_2 est plus efficace que celui obtenu avec l'information mutuelle en considérant la connectivité et les artefacts locaux, alors que le noyau issu de l'information mutuelle est légèrement meilleur en ce qui concerne BV/TV et le terme d'attache aux données.

Conclusion

La détermination du noyau est un enjeu majeur dans la restauration d'images. Dans ce chapitre, nous avons comparé deux méthodes de détermination du noyau flou sur la base de l'optimisation de la norme L_2 ou de l'information mutuelle. Des résultats similaires sont obtenus avec les deux noyaux dérivés de la norme L_2 ou l'information mutuelle. Le noyau de la norme L_2 préserve une meilleure densité de connectivité, et celui de l'information mutuelle a une meilleure performance pour le rapport BV/TV. Une amélioration possible pour notre méthode de descente du gradient appliqué à l'information mutuelle est de considérer l'information locale ([Hermosillo et al., 2002](#)). Hermosillo *et al.* a proposé une méthode d'information mutuelle locale, qui va générer un noyau variant spatialement pour améliorer les résultats. Ces résultats doivent être confirmés sur d'autres ROI et d'autres images HR-pQCT.

De plus, la super résolution d'images basées sur une régularisation TV montre que les noyaux obtenus améliorent la qualité des images HR-pQCT, mais c'est encore loin de nos attentes. Nous envisageons d'explorer des a priori différents pour résoudre cette tâche de super résolution. Dans le chapitre suivant, nous intégrons a priori sur l'image de super résolution afin d'améliorer le contraste des images.

Chapitre 4 : Super résolution par optimisation nonconvexe et nonlisse

Une image avec un bon contraste favorise la tâche de segmentation, en particulier pour l'extraction de la micro-architecture trabéculaire à partir d'images CT avec niveaux de gris. La vérité terrain obtenue à partir de μ -CT a une nature quasi binaire et son histogramme suit une distribution bimodale. Cependant, cette caractéristique est perdue dans les images à basse résolution générées à partir de HR-pQCT. Fig.1 illustre l'histogramme des images à haute et basse résolution.

Nous essayons de récupérer la bimodalité des histogrammes de super résolution, et un bon candidat est d'inclure un double puits de potentiel dans le terme de régularisation malgré sa non convexité. Basé sur le travail précédent (Li et al., 2017a), nous proposons de combiner une contrainte non convexe à double puits avec une fonction de régularisation TV non lisse pour améliorer la bimodalité des images résultantes.

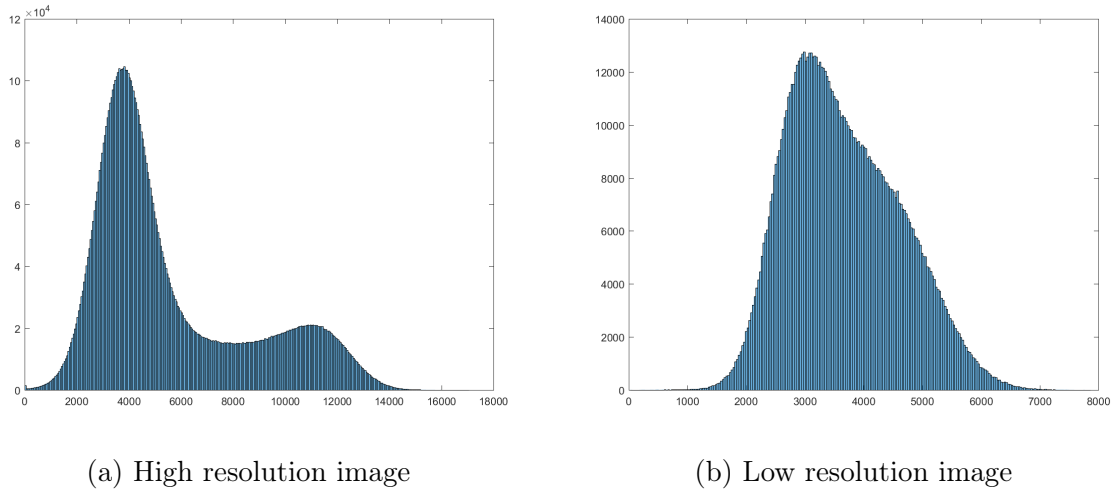


Figure 1: Histogramme typique haute résolution (taille de voxel $41\mu\text{m}$) et basse résolution (taille de voxel $82\mu\text{m}$) d'images CT. Suivant (a), s_1 et s_2 correspondent aux positions des deux pics dans l'histogramme de l'image haute résolution ($s_1 = 3750$, $s_2 = 11075$).

Méthodes

La fonction régularisée pour résoudre le problème de super résolution est écrite comme suit :

$$\mathcal{J}(\mathbf{f}) = \|\mathbf{A}\mathbf{f} - \mathbf{g}\|_2^2 + \mu\mathcal{R}(\mathbf{f}) \quad (5)$$

Dans le chapitre 3, $\mathcal{R}(\cdot)$ est la régularisation TV: $\sum_i \|\mathcal{D}_i \mathbf{f}\|_2$, où $\mathcal{D}_i$ prend le gradient du voxel de $\mathbf{f}$ à l'indice i . Dans ce chapitre, $\mathcal{R}(\cdot)$ inclut la régularisation TV et la fonction double puits. Le potentiel de double puits est écrit comme suit

$$\mathcal{W}(\mathbf{f}_i) = (\mathbf{f}_i - s_1)^2(\mathbf{f}_i - s_2)^2 \quad (6)$$

où $s_1, s_2 \in \mathbb{R}^+$ sont les positions des pics de l'histogramme bimodal, ils sont fixés empiriquement, et i est l'indice du voxel. La fonctionnelle est formulée ainsi:

$$\mathcal{J}(\mathbf{f}) = \|\mathcal{A}\mathbf{f} - \mathbf{g}\|_2^2 + \mu_1 \sum_i \|\mathcal{D}_i \mathbf{f}\|_2 + \mu_2 \sum_i \mathcal{W}(\mathbf{f}_i) \quad (7)$$

où μ_1, μ_2 sont les paramètres de régularisation. Trois méthodes sont proposées pour optimiser cette fonctionnelle.

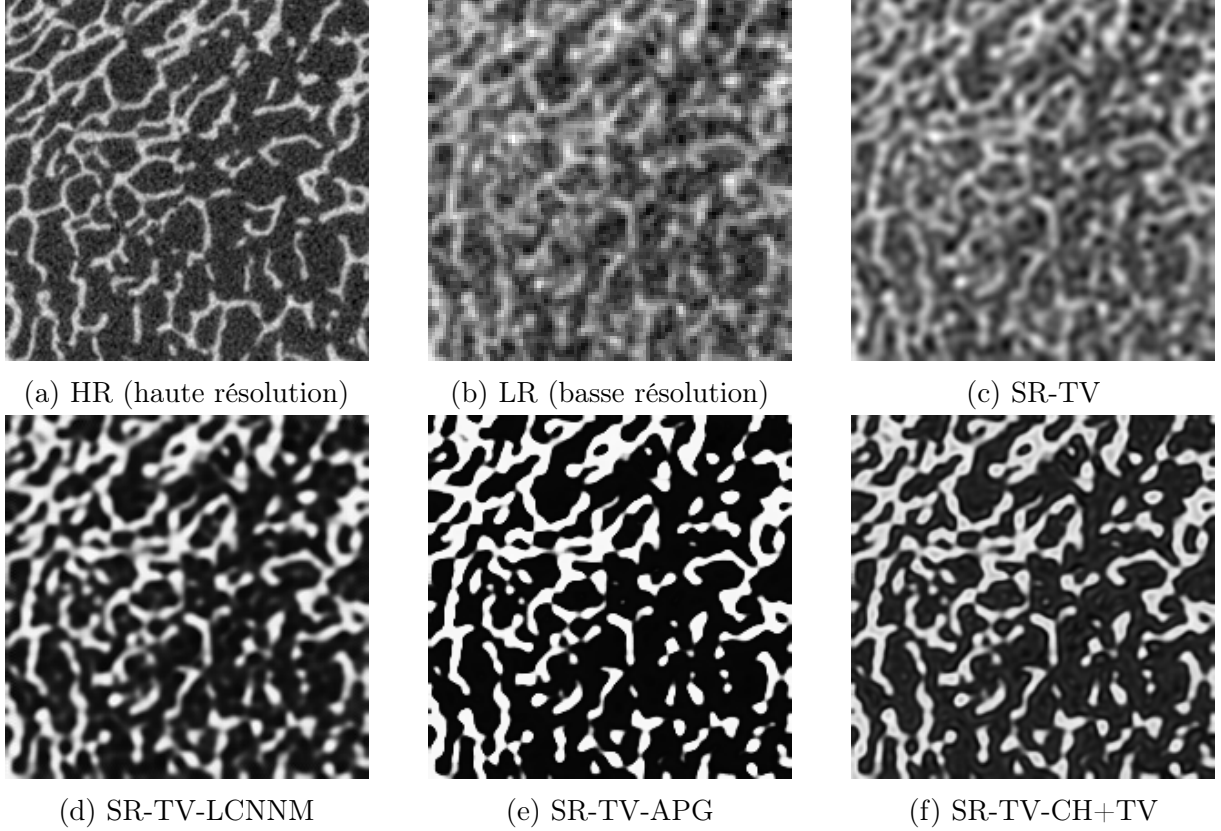


Figure 2: Tranches 2D d'un morceau test: HR, LR et SR désignent les images haute résolution, basse résolution et super résolution. SR-TV est l'image super résolution restaurée par TV. (d)-(f) représentent des images reconstruites par LCNNM, APG et CH+TV initialisées l'image de super résolution obtenue avec TV, SR-TV.

Expériences numériques

Les tests ont été effectués sur deux images d'os trabéculaire expérimentales en basse résolution obtenues à partir de scanners HR-pQCT (Boutroy et al., 2005) avec une taille de voxel de $82\mu\text{m}$. Deux autres images à haute résolution ont également été obtenues sur les mêmes échantillons d'os à partir avec la $\mu\text{-CT}$ avec une taille de voxel de $41\mu\text{m}$ après recalage. Le recalage a été effectué avec un logiciel open source Elastix (Klein et al., 2010, Shamonin et al., 2014) avec l'interpolation B-spline et la méthode stochastique de descente de gradient. L'opérateur de floutage que nous avons appliqué dans ce travail a été déterminé avec la minimisation de la norme L_2 , comme expliqué dans le chapitre

3. Une première approche de restauration basée sur la régularisation TV a été utilisée comme point de départ pour les trois méthodes de minimisation. Cette méthode a été développée dans un travail précédent et est détaillée dans (Toma et al., 2014a, Toma et al., 2015). Les fonctions de régularisation sont non convexes et les résultats obtenus avec des images de départ non régularisées sont très médiocres.

Résultats

Fig.2 affiche en 2D les images à haute et basse résolution, l'image initiale, les images obtenues avec TV et les trois approches proposées. Les images super résolution obtenues avec les trois algorithmes ont un bon contraste et un indice de Dice amélioré. Cependant, leur connectivité est dégradée.

Conclusion

Les images reconstruites par les trois approches ont un contraste amélioré dû au double puits et une structure lisse du fait de la TV. En raison de l'absence de convexité de la régularisation fonctionnelle, des structures similaires sont obtenues avec les différents algorithmes, mais leur paramètres morphologiques et topologiques sont différents. Nous avons montré qu'un arrêt précoce des algorithmes peut être utile pour obtenir les meilleurs résultats. La méthode CH+TV est considérée comme la meilleure méthode parmi les trois algorithmes. Pourtant, les méthodes proposées ne préservent pas la connectivité de la structure.

Les résultats montrent que les images super résolution ont perdu beaucoup de détails par rapport à la vérité terrain. Nous envisageons d'ajouter plus d'informations structurales sur les travées osseuses dans le but de récupérer plus de détails. L'apprentissage par dictionnaire est un bon candidat car il permet d'extraire des images caractéristiques d'une région de perception plus grande que les régularisation TV ou CH+TV. Des approches d'apprentissage par dictionnaire ont été introduites dans Sect. 2.2.4. Le principe est d'apprendre deux dictionnaires (Φ^g , Φ^f) à partir d'un ensemble de paires d'images basse et haute résolution pendant la phase de l'apprentissage. Ensuite, au stade de la super résolution, les images à basse résolution peuvent être exprimées par un certain nombre d'atomes dans le dictionnaire Φ^g . Il faut ensuite utiliser leurs atomes correspondants dans le dictionnaire Φ^f pour reconstruire les images à haute résolution. Nous allons introduire l'application de l'apprentissage du dictionnaire dans notre problème dans le prochain chapitre.

Chapitre 5 : Super résolution avec l'apprentissage dictionnaire

Nous étudions une approche basée sur la méthode Semi coupled dictionary learning (Wang et al., 2012)(SCDL) pour augmenter la résolution. Pour autant que nous le sachions, il s'agit du premier travail d'investigation de SCDL pour des images expérimentales de tomodynamométrie avec une architecture complexe. De plus, l'application de cette approche en 3D et la comparaison avec d'autres méthodes basées sur la TV ou la TV couplée à un double puits potentiel est originale. Le traitement des images 3D ajoute la complexité du processus par rapport au cas 2D étudié dans le travail original (Wang et al., 2012) en raison de l'augmentation des besoins en mémoire et en temps de calcul. Toutefois, la caractérisation de la connectivité et des propriétés topologiques de l'os nécessite des images 3D. Afin de gagner du temps de calcul et de gérer l'anisotropie osseuse, nous proposons une stratégie 2.5D pour apprendre les dictionnaires haute et basse résolution sur les directions coronale/sagittale/axiale séparément.

Méthodes

La méthode SCDL (Wang et al., 2012) apprend deux dictionnaires, un dictionnaire pour les images haute résolution, désigné comme Φ^f , l'autre dictionnaire pour les images en basse résolution, désigné comme Φ^g . Désignons par P and Q la taille du patchs et le nombre de patchs respectivement: pour $1 \leq i \leq Q$, $\mathbf{f}_i \in \mathbb{R}^{P \times 1}$ est un patch appartenant au dataset des images hautes résolution, $\mathbf{g}_i \in \mathbb{R}^{P \times 1}$ est un patch en basse résolution après interpolation. Le nombre des atomes dans Φ^f et Φ^g sont K^f et K^g , donc $\Phi^f \in \mathbb{R}^{P \times K^f}$, $\Phi^g \in \mathbb{R}^{P \times K^g}$, $\alpha_i^f \in \mathbb{R}^{K^f \times 1}$ et $\alpha_i^g \in \mathbb{R}^{K^g \times 1}$ sont les représentations sparses des patchs $\mathbf{f}_i$ et $\mathbf{g}_i$, $\alpha^f \in \mathbb{R}^{K^f \times Q}$, $\alpha^g \in \mathbb{R}^{K^g \times Q}$.

Même si $\mathbf{g}$ (après interpolation) et $\mathbf{f}$ peuvent être représentés de façon sparse par α^g et α^f avec Φ^g et Φ^f respectivement, un opérateur de correspondance associe α^g et α^f . Dans (Wang et al., 2012), deux matrices de correspondance $\mathbf{W}^f$ et $\mathbf{W}^g$ sont introduites pour que $\alpha^g = \mathbf{W}^f \alpha^f$ et $\alpha^f = \mathbf{W}^g \alpha^g$, donc $K^f = K^g$. Le fonctionnelle objective pour apprendre les deux dictionnaires est écrite comme :

$$\begin{aligned} \min_{\Phi^g, \Phi^f, \mathbf{W}^g, \mathbf{W}^f, \alpha^f, \alpha^g} & \|\mathbf{g} - \Phi^g \alpha^g\|_F^2 + \|\mathbf{f} - \Phi^f \alpha^f\|_F^2 + \gamma \|\alpha^f - \mathbf{W}^g \alpha^g\|_F^2 + \gamma \|\alpha^g - \mathbf{W}^f \alpha^f\|_F^2 \\ & + \lambda^W \|\mathbf{W}^g\|_F^2 + \lambda^W \|\mathbf{W}^f\|_F^2 + \lambda^f \|\alpha^f\|_1 + \lambda^g \|\alpha^g\|_1 \\ \text{s.t.} & \|\mathbf{d}_i^g\|_2 \leq 1, \|\mathbf{d}_i^f\|_2 \leq 1, \forall i \end{aligned} \quad (8)$$

où $\|\cdot\|_F$ est la norme de Frobenius et i est l'indice de colonne. Les variables dans Eq.8 sont mises à jour alternativement.

Quand les deux dictionnaires sont appris, nous observons que certains atomes ne sont pas utiles pour la reconstruction des images de super résolution, parce qu'ils semblent très bûités. Nous avons choisi un critère qui permet d'évaluer la similarité avec du bruit, et par conséquent, certains atomes sont retirés des dictionnaires.

Ensuite, le problème de super résolution peut être résolu par:

$$\begin{aligned} \min_{\alpha_i^f, \alpha_i^g, f_i} & \|g_i - \Phi^g \alpha_i^g\|_F^2 + \|f_i - \Phi^f \alpha_i^f\|_F^2 + \\ & \gamma' \|\alpha_i^f - W^g \alpha_i^g\|_F^2 + \gamma' \|\alpha_i^g - W^f \alpha_i^f\|_F^2 + \\ & \lambda^g \|\alpha_i^g\|_1 + \lambda^f \|\alpha_i^f\|_1 \end{aligned} \quad (9)$$

Trois dictionnaires d'images à haute résolution sont appris pour les problèmes de super résolution 3D. Chaque dictionnaire présente des images caractéristiques dans une direction. Les trois directions coronaire/sagittale/axiale ont été prises en compte. Fig.3 illustre le topologie de la méthode 2.5D pour reconstruire des images 3D de super résolution.

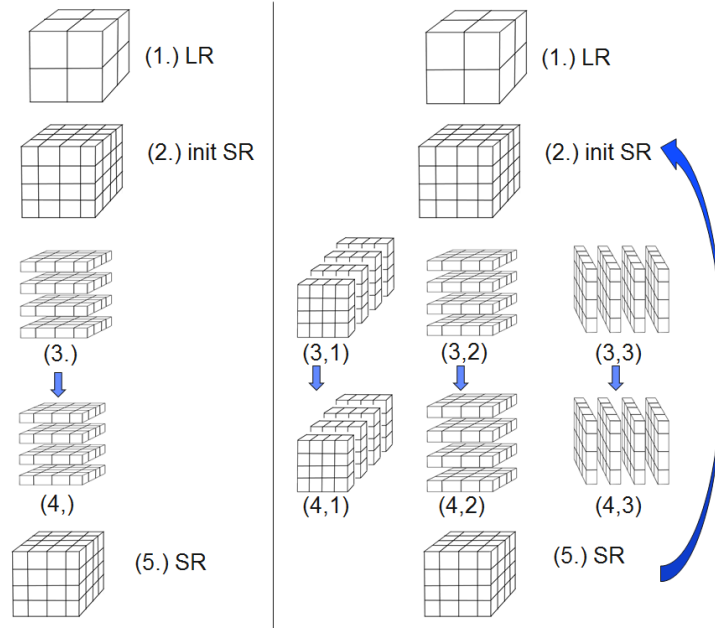


Figure 3: La partie de gauche est le schéma de la méthode 2D SCDL pour la reconstruction d'image de super résolution 3D; la partie de droite est le schéma de la méthode 2.5D SCDL pour la reconstruction d'image de super résolution 3D. (1.)LR est l'image à faible résolution; (2.)init SR est la super résolution initiale en 3D; (3.)(3.1)(3.1)(3.2)(3.3.3)(3.3) sont des tranches 2D coupées à partir du volume 3D dans des directions différentes (Notez que les volumes (3.) et (3.2) sont coupés horizontalement); (4.)(4.1)(4.2)(4.3)(4.3) sont des images de super résolution restaurées avec SCDL en dimension 2D; (5.) à gauche est représenté un volume empilé avec des tranches obtenues précédemment (4.); (5.) à droite prend la moyenne de (4.1)(4.2)(4.2)(4.3).

Expériences numériques

Les données d'imagerie ont été recueillies pour un fantôme de structure osseuse (Burghardt et al., 2013). Le fantôme était composé de sections de radius distal cadavérique et de tibia (1 cm d'épaisseur) noyées dans un cylindre de 8 cm de diamètre composé de polyméthacrylate de méthyle (PMMA) et de résine de polyéthylène. Le fantôme a été scanné en utilisant les paramètres du protocole typique *in vivo* par HR-pQCT (XtremeCT, Scanco Medical

AG), qui génère des reconstructions avec des voxels de $82\mu\text{m}$ isotropes. Les données de référence ont été collectées sur le système de laboratoire $\mu\text{-CT}$ (Scanco Medical $\mu\text{-CT}$ 100), qui a fourni des reconstructions à haute résolution, avec $24\mu\text{m}$ de voxels isotropes, après recalage et interpolation B-spline, la taille de voxel est égale à $41\mu\text{m}$.

La taille du voxel des images à basse et haute résolution est respectivement de $82\mu\text{m}$ et $41\mu\text{m}$. La taille des images HR-pQCT est de $420 \times 510 \times 110$, et celle des images CT est de $840 \times 1020 \times 220$. Sept échantillons d'os ont été utilisés dans ce chapitre, trois tibia et cinq radius.

Résultats

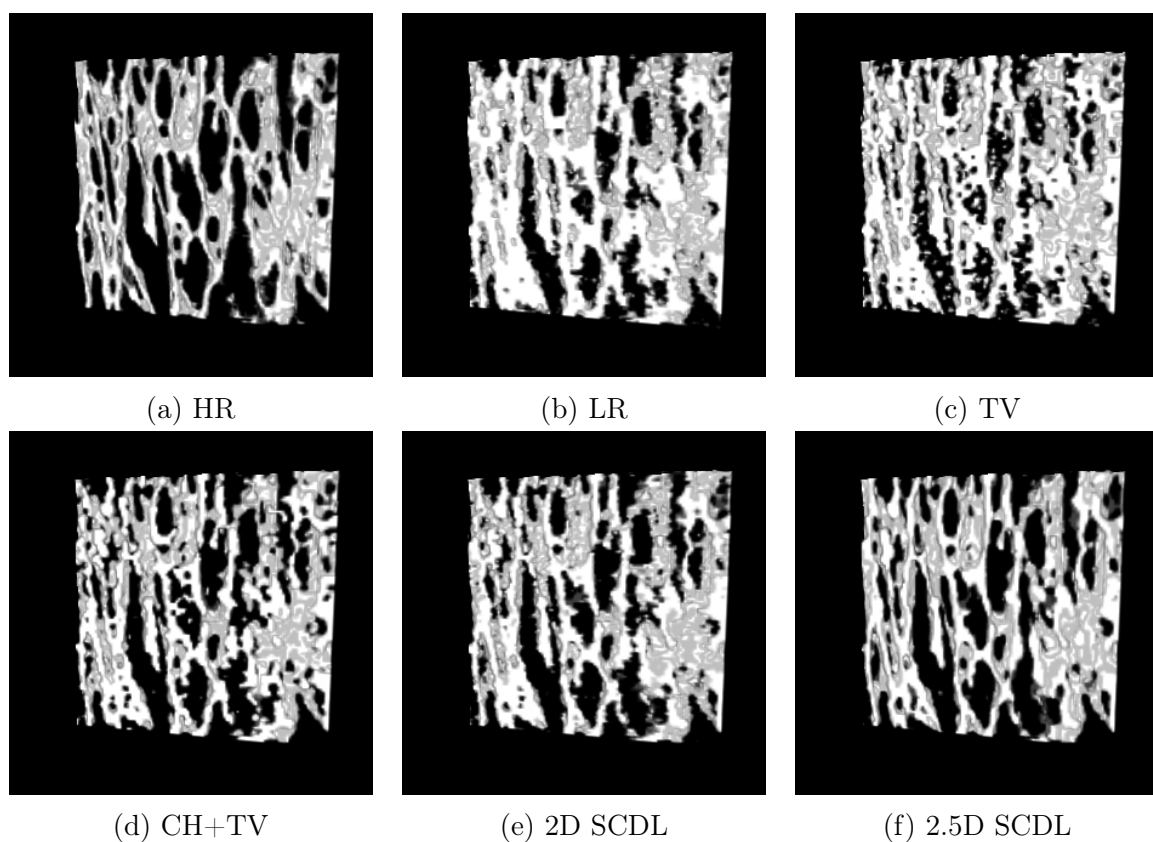


Figure 4: Illustration 3D d'images limitées à une pile de 10 tranches, sur la première ligne, de gauche à droite, haute résolution, basse résolution, super résolution obtenue avec le TV; sur la deuxième ligne, respectivement CH+TV, 2D SCDL et 2.5D SCDL. L'image affichée est un crop qui vient du set d'apprentissage et est coupé dans l'échantillon A_{tibia} , les tailles des patches 2D et 2.5D SCDL sont 12^2 .

Afin de comparer visuellement les différentes méthodes de super-résolution appliquées aux images CT expérimentales, Fig.4 montre 6 piles de l'échantillon coupé à partir de volumes 3D obtenus avec les différentes méthodes.

L'image obtenue par la régularisation TV améliore la qualité des images par rapport à la basse résolution, mais elle est corrompue par le bruit. L'image obtenue avec CH+TV semble diminuer les effets de bruit par rapport à l'images obtenue avec TV, mais sa micro-architecture n'est pas aussi mince que celle de la vérité terrain. Avec 2D SCDL, l'image

est moins bruitée, mais sa structure n'est pas lisse. 2.5D SCDL récupère une plus fine structure avec une meilleure texture.

Conclusion

Dans ce chapitre, nous avons développé une méthode SCDL pour résoudre le problème de super résolution 3D HR-pQCT. La méthode du dictionnaire a été mise en œuvre sur des tranches dans les trois directions spatiales. Les résultats de super résolution sont comparés à ceux obtenus avec TV ou TV combiné avec le potentiel de Cahn-Hilliard. Les images restaurées par 2.5D SCDL sont meilleures que celles obtenus avec les autres méthodes. Nous aimerions appliquer cette approche à des échantillons expérimentaux plus importants afin de valider la conclusion et d'étendre l'analyse morphologique.

Chapitre 6 : Revue de l'apprentissage profond pour la super résolution

Récemment, les approches d'apprentissage profond ont connu un grand succès dans le traitement d'image avec des tâches telles que la classification ([Krizhevsky et al., 2012](#)), la segmentation ([Ronneberger et al., 2015](#)), le débruitage ([Zhang et al., 2017](#)). En fait, l'apprentissage par dictionnaire et l'apprentissage profond font tous deux partie de l'apprentissage machine. Cependant, les approches d'apprentissage profond présentent deux avantages principaux qui les distinguent des autres approches : un calcul parallèle très développé et une puissante capacité de représentation.

(1) L'apprentissage profond permet d'apprendre les caractéristiques de haut niveau des données d'une manière similaire au traitement visuel dans le cerveau humain en utilisant plusieurs couches de neurones avec des non linéarités ([LeCun et al., 2015](#)). Une grande quantité de paramètres est utilisée dans les réseaux d'apprentissage profond afin de révéler des informations implicites et d'approximer l'opérateur inverse du noyau de dégradation. Pourtant, l'efficacité de l'approche d'apprentissage profond est liée à la quantité de données. Une grande quantité de données peut améliorer les performances de l'apprentissage profond, au contraire, un nombre limité de données limite ses performances. Par ailleurs, l'interprétation du réseau de neurones d'apprentissage profond reste un défi.

(2) Le calcul parallèle a été beaucoup développé dans le cadre de l'apprentissage profond, comme avec Tensorflow ([Abadi et al., 2016](#)), MXNet ([Chen et al., 2015](#)), Caffe ([Jia et al., 2014](#)). L'utilisateur peut bénéficier directement du calcul parallèle à grande vitesse sans connaître l'architecture GPU et la programmation GPU de bas niveau.

Dans ce chapitre, tout d'abord, nous introduisons des réseaux d'apprentissage en profondeur pour les problèmes de super résolution. Évoluant du super resolution convolutional neuron network au réseau profond tel que le very deep residual network pour la super résolution (VDSR), du réseau récursif au réseau résiduel récursif, du réseau résiduel au réseau densément connecté, l'architecture d'apprentissage profond tend à être de plus en plus robuste et plus fortement connectée.

Ensuite, nous donnons quelques exemples de l'application de l'apprentissage en profondeur pour le traitement d'images médicales. Les méthodes d'apprentissage approfondi ont un vaste champ d'application dans les tâches de traitement d'images médicales ([Litjens et al., 2017](#)), comme la classification, la détection, la segmentation, etc. Les recherches en traitement d'images médicales ont été influencées par les méthodes d'apprentissage profond, et proposent maintenant de nouvelles architectures permettant d'augmenter la performance du réseau et de faciliter l'analyse et le suivi des recherches.

Plus de détails sont disponibles dans le chapitre 6.

Chapitre 7 : Super résolution avec machine learning

Dans le chapitre 5, nous avons discuté de l'application de SCDL pour les images d'os trabéculaire. Et les résultats montrent que SCDL améliore effectivement la qualité des images. Cependant, l'un des défis de l'utilisation des méthodes d'apprentissage des dictionnaires pour la super résolution est de savoir comment coupler efficacement les dictionnaires haute et basse résolution. Différentes fonctions ont été discutées dans (Yang et al., 2010a, Yang et al., 2012, Wang et al., 2012, He et al., 2013). Pourtant, ces fonctions sont supposées être linéaires, ce qui peut ne pas être suffisant et on peut envisager des fonctions plus complexes. Le réseau d'apprentissage profond est un bon candidat pour décrire des fonctions non linéaires.

Dans ce chapitre, nous proposons un réseau FSRCNN-Réseau résiduel pour résoudre le problème joint super-résolution/segmentation pour les images HR-pQCT. Même si FSRCNN n'est pas la meilleure architecture pour les problèmes de super résolution, il a moins de paramètres et fournit des résultats acceptables. Le réseau résiduel affine la segmentation. Le réseau proposé est d'abord appliqué pour résoudre un problème de super résolution avec le facteur de sur-échantillonnage $r = 2$ et les résultats sont comparés avec les méthodes précédemment étudiées. Puis une modification mineure est introduite pour que le réseau résolve un problème de super résolution avec le facteur $r = 3.42$.

Méthode

Avant d'introduire l'architecture, nous rappelons d'abord quelques notations représentant différents types d'images. Soit $\mathbf{g} \in \mathbb{R}^N$ les images 3D HR-pQCT basse résolution, $\mathbf{f}^* \in \mathbb{R}^{N'}$ représente la vérité terrain 3D haute résolution obtenue à partir de μ -CT, $\mathbf{f}_b^* \in \mathbb{R}^{N'}$ est l'image binaire de $\mathbf{f}^*$. $\mathbf{f}$ et $\mathbf{f}_b \in \mathbb{R}^{N'}$ sont les images super résolution de niveau de gris et les images binaires, respectivement. $\mathbf{f}_s \in \mathbb{R}^{N'}$ est la sortie du réseau, prédisant la probabilité que les voxels appartiennent au tissu osseux. Le facteur d'échantillonnage isotrope est de 2, donc $N' = 8N$ pour les images 3D.

L'architecture du réseau proposé se compose de FSRCNN et d'un réseau résiduel, comme le montre la Fig.5. Les tailles des filtres sont précisées dans le schéma. Mise à part, la dernière fonction d'activation adjacente à la couche de sortie qui est sigmoïde, toutes les fonctions d'activation sont de type PReLU. Plus de détails sont présentés au chapitre 7. Les fonctions de perte du réseau sont définies comme suit:

$$\begin{aligned}\mathcal{L}_1(\mathbf{f}^*, \mathbf{f}) &= \frac{1}{MN'} \sum_i^M \|\mathbf{f}_i^* - \mathbf{f}_i\|_2 \\ \mathcal{L}_2(\mathbf{f}_b^*, \mathbf{f}_s) &= - \frac{1}{MN'} \sum_i^M \sum_j^{N'} (\mathbf{f}_{b_i,j}^* \log(\mathbf{f}_{s_i,j}) + (1 - \mathbf{f}_{b_i,j}^*) \log(1 - \mathbf{f}_{s_i,j})) \\ \mathcal{L}(\mathbf{f}^*, \mathbf{f}_b^*, \mathbf{f}, \mathbf{f}_s) &= \mathcal{L}_1(\mathbf{f}^*, \mathbf{f}) + \mathcal{L}_2(\mathbf{f}_b^*, \mathbf{f}_s)\end{aligned}\tag{10}$$

$\mathcal{L}_1$ tient compte de l'estimation de la super résolution à partir des niveaux de gris avec la norme L_2 . $\mathcal{L}_2$ affine les effets de segmentation avec un terme de "cross entropy". Plus de détails sont présentés dans le chapitre 7.

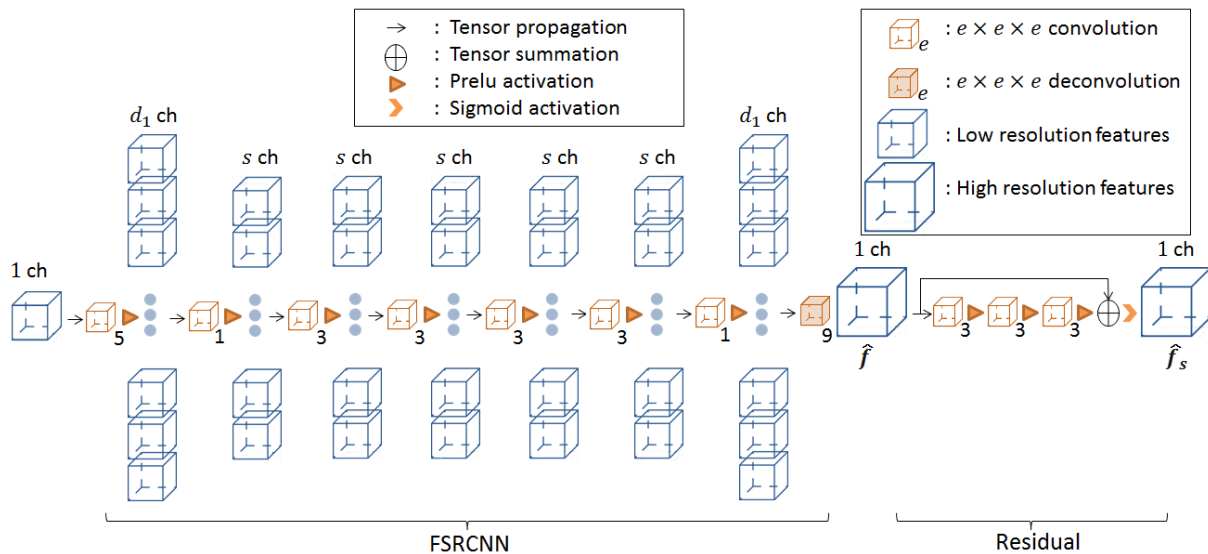


Figure 5: Illustration de l’architecture FSRCNN-Residual network. La première partie FSRCNN(Dong et al., 2016) traite le problème super résolution, la dernière partie Residual(Kim et al., 2016a) sert à la segmentation. Les symboles $\hat{\mathbf{f}}$ et $\hat{\mathbf{f}}_s$ représentent les images super résolution et la sortie du réseau, respectivement. Les cubes bleus représente les features, les cubes oranges représentent les filtres. Les numéros au dessus des cubes bleus sont les nombres des filtres (d_1, s, d_2), d_2 est le nombre de filtre dans le réseau Residual sub-network, celui qui n’est pas représenté dans cette figure. Les numéros à coté des cubes oranges sont les longueurs du bord des filtres 3D. Par exemple, quand les longueurs du bord égale à 5, la taille de ce filtre est $5 \times 5 \times 5$.

Expériences numériques

Notre base de données est composée de 13 échantillons de cadavres dont 7 radius et 6 tibias. La taille du patch est de $14 \times 14 \times 14$ pour les images basse résolution, la taille de la paire haute résolution correspondante est de $20 \times 20 \times 20$. Les nombres de filtres sont $d_1 = 64$, $s = 24$, le nombre de filtres est de 32 pour les trois couches du réseau résiduel. Le taux d'apprentissage est de 10^{-4} . Le réseau est optimisé avec la méthode de descente de gradient stochastique. Le code est implémenté en Tensorflow et est lancé avec le NVIDIA Titan Xp.

Résultats

Fig.6 compare les différentes performances des méthodes distinctes de super résolution et montre que le framework FSRCNN-Residual permet de reconstruire des images de super résolution avec une structure fine basée sur des images basse résolution. Nous avons également calculé les paramètres osseux relatifs pour comparer quantitativement différentes approches de super résolution. Bien que l'indice de connectivité Conn.D ne soit pas amélioré, les indices Dice et BV/TV obtenus surpassent toutes les autres approches.

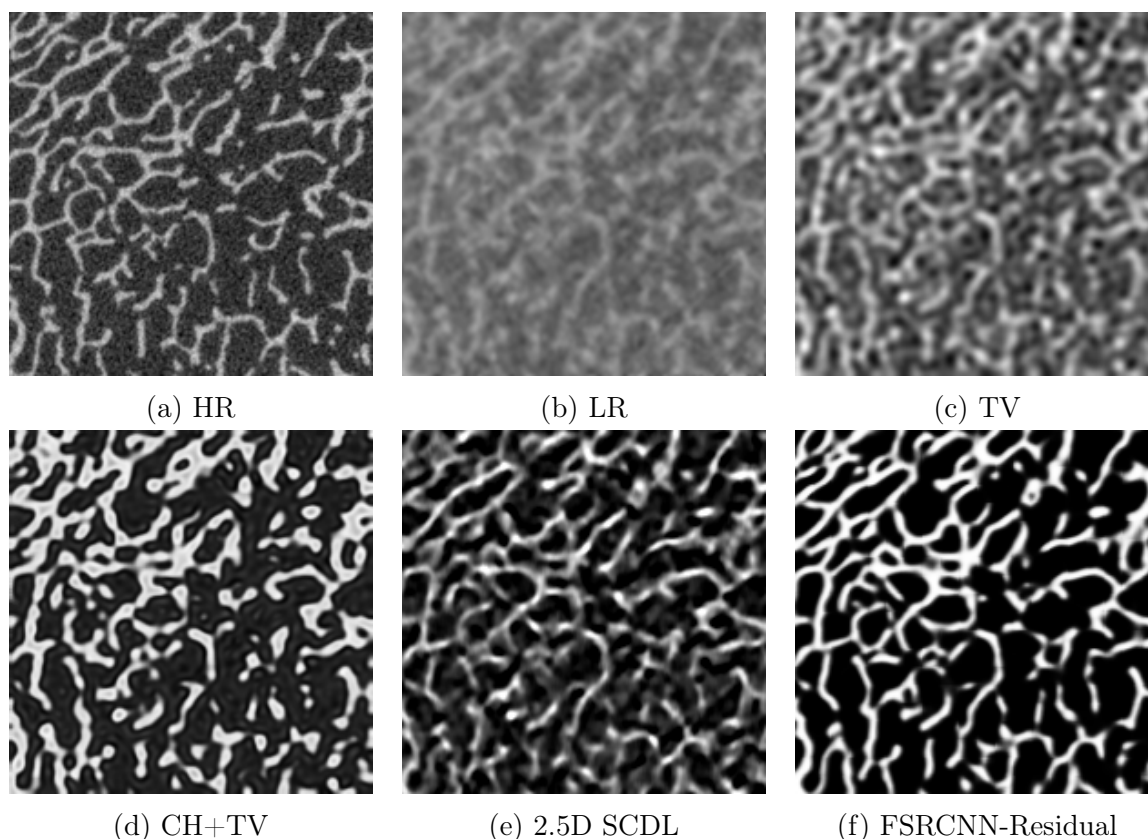


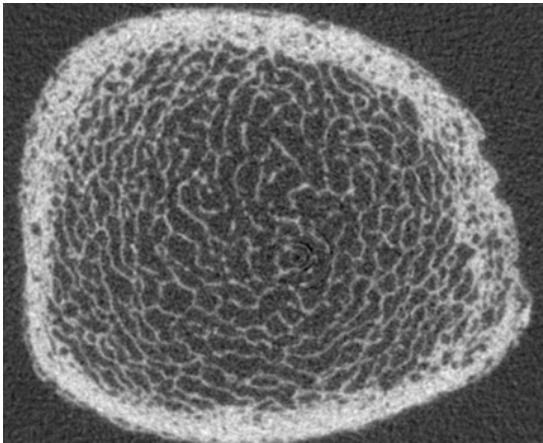
Figure 6: Illustration des images 2D coupées de volumes 3D en niveau de gris. Le volume vient de l'ensemble d'apprentissage. Sur la première ligne, de gauche à droite, haute résolution, basse résolution, super résolution, image TV super résolution ; sur la deuxième ligne, les images obtenues avec CH+TV, 2.5D SCDL et l'approche proposée, respectivement.

Super résolution avec facteur de sur-échantillonnage 3,42

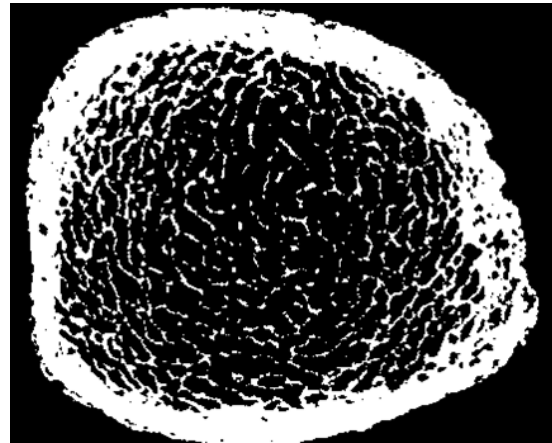
La convolution est un opérateur qui permet de produire des images de taille flexible. Nous proposons donc de modifier le réseau FSRCNN-Résiduel pour résoudre le problème de super résolution de la taille de voxel $82\mu\text{m}$ à $24\mu\text{m}$.

L'architecture est similaire à celle proposée dans la section précédente, qui se compose du FSRCNN et du réseau Résiduel. FSRCNN sert à la super résolution des problèmes, le réseau résiduel affine la segmentation. Une grande différence est que la taille du filtre de la couche de déconvolution est modifiée à 8. Les détails sont présentés au chapitre 7.

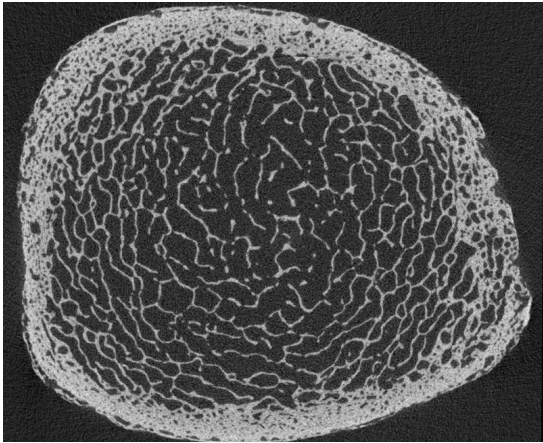
L'échantillon dans la Fig. 7 est un échantillon de tibia prélevé dans l'ensemble test. L'image de super résolution sur cette figure montre que le réseau conserve son efficacité sur les images dans l'ensemble test. L'architecture de l'image de super résolution a été récupérée à partir d'images à basse résolution et l'épaisseur de trabeculea est visuellement comparable à celle des images à haute résolution. De plus, les régions corticales ont également été améliorées. Pourtant, les régions corticales sur les images super résolution ne sont pas aussi poreuses que sur les images haute résolution. Cela peut s'expliquer par le fait que la quantité de patches contenant des os corticaux est limitée dans l'ensemble d'apprentissage, tandis que les variations sur les régions corticales peuvent s'estomper en raison de la normalisation sur l'ensemble des images.



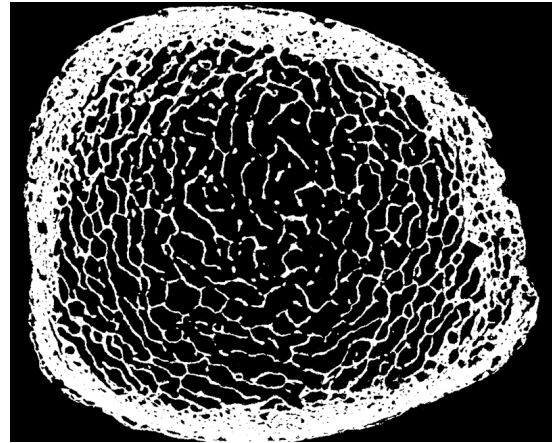
(a) image basse résolution de niveau gris



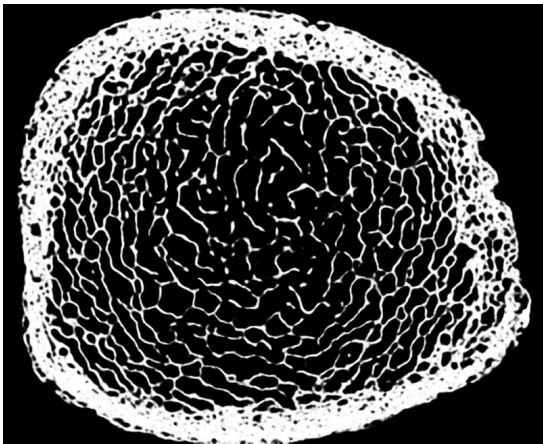
(b) image basse résolution en binaire



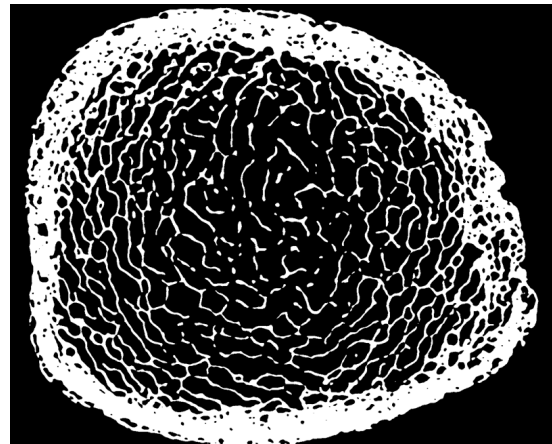
(c) image haute résolution de niveau gris



(d) image haute résolution en binaire



(e) image super résolution de niveau gris



(f) image super résolution en binaire

Figure 7: Illustration de tranches 2D d'échantillons entiers en 3D à la taille voxel $24\mu\text{m}$. L'échantillon testé est un tibia dans l'ensemble de test. Cette figure est utilisée pour montrer la généralisation du réseau. Les images de la colonne de gauche sont des images de niveau de gris, de haut en bas : basse résolution, haute résolution, image super résolution reconstruite par notre méthode proposée. Toutes les images de la colonne de droite sont des images binaires. Ils sont segmentés avec Otsu, de haut en bas : basse résolution, haute résolution et image super résolution reconstruite par notre méthode.

Conclusion

Dans ce chapitre, nous avons proposé un réseau FSRCNN-Résiduel pour résoudre un problème de super résolution/segmentation joint. Les deux fonctions de perte dans le réseau injectent des a priori pour les images de niveau de gris et les images binaires. Nous proposons d'utiliser le réseau FSRCNN-Residual pour résoudre un problème de super résolution/ segmentation. Le facteur de sur-échantillonnage est de 2 et 3,42. Nos résultats expérimentaux pour le facteur de sur-échantillonnage égal à 2 montrent que le réseau proposé surpasse les approches étudiées précédemment et qu'il est très prometteur pour la recherche future. Comme cette approche nous permet de traiter des échantillons entiers, l'étape suivante est l'analyse morphologique.

Conclusions et perspectives

Conclusions

Dans cette thèse, nous tentons de résoudre un problème de super résolution/segmentation joint pour des images expérimentales 3D HR-pQCT en micro architecture osseuse trabéculaire. Les images HR-pQCT sont obtenues à partir de mesures *in vivo* et la taille du voxel est de $82\mu\text{m}$ dans cette thèse. Les images de la vérité terrain sont collectées à partir d'images *ex vivo* $\mu\text{-CT}$ et la taille du voxel est de $24\mu\text{m}$. Notre objectif est d'améliorer la résolution des images HR-pQCT afin que les paramètres osseux extraits des images super résolution segmentées soient proches de ceux obtenus avec les mesures d'images $\mu\text{-CT}$. Nous avons essayé différentes approches, passant d'approches variationnelle à l'apprentissage par dictionnaires, puis à des méthodes deep learning. Notre expériences numériques montrent que la qualité de l'image s'est progressivement améliorée.

Perspectives

Une façon de préserver la connectivité est d'introduire une perte supplémentaire qui tient compte des informations locales du voisin. Le réseau actuel a deux pertes, mais les deux calculent la précision de l'élément d'estimation. Les a priori concernant la similarité des voisins peuvent améliorer les performances sur la connectivité.

De plus, en s'inspirant du réseau de neurones à contraintes anatomiques, l'amélioration de la similarité de l'estimation et des images à haute résolution dans un espace de basses dimensions peut aider à préserver la connectivité. Le défi consiste à définir l'espace de basse dimension qui décrit la structure topologique.

General introduction

Osteoporosis is a disease characterized by loss of bone mass and degradation of bone microarchitecture. Osteoporosis is one of the most common bone diseases among the elderly. Although osteoporosis is not a fatal disease, the fractures it causes can lead to serious complications (damage to vessels and nerves, infections, stiffness), sometimes accompanied by threats of death.

Bone mineral density (BMD) is a criterion for diagnosing osteoporosis. But only BMD is not sufficient to quantify bone quality. Micro architecture plays an important role in the diagnosis of osteoporosis. Since it provides parameters related to bone strength, the precise characterization of the bone microstructure is thus crucial for diagnosis.

High resolution peripheral computed tomography (HR-pQCT) is currently one of the best CT techniques for studying bone microarchitecture in 3D *in vivo*. It scans bone microarchitecture *in vivo* at peripheral human sites, typically with a voxel size of $82\mu\text{m}$, the spatial resolution ranging from $130\mu\text{m}$ - $150\mu\text{m}$. However, such spatial resolution is still insufficient because the average size of trabeculae is around $150\mu\text{m}$. Therefore, it is difficult to have an accurate estimation of bone parameters, which are calculated after bone segmentation. Micro-CT (μ -CT) can provide images at much higher spatial resolution, but is limited to *ex vivo* studies.

In this thesis, our aim is to solve a joint super resolution/segmentation problem for HR-pQCT images and to obtain a good structural characterization of the bone with super resolution images. In all our work, we are working on experimental images. HR-pQCT images with voxel size μm will be considered as low resolution images and μ -CT images of the same specimen at a voxel size of $24\mu\text{m}$, will be considered as high resolution images and used as the ground truth for the bone microarchitecture after registration.

This manuscript is organized in seven chapters, the first two being devoted to the presentation of the medical context and background in inverse problems and the five next, to our own contributions.

The chapter 1 is devoted to the presentation of the medical context: we detail the importance of the trabecular bone micro architecture and the different tomodensitometry CT techniques, then, we give the definitions of the usual bone parameters to quantify the quality of the trabecular bone. Next, we discuss the advantages and disadvantages of HR-pQCT images and show our motivations for this work.

The chapter 2 summarizes the state of the art of super resolution methodologies and present some mathematical notions on optimization methods. The methods to solve a super resolution problem can be classified into 4 types: (1) spatial interpolation, (2) frequency domain image processing, (3) function optimization with regularization terms, (4) machine learning. In this thesis, we focus on the last two methodologies. After the introduction of the state of the art on super resolution problems, we introduce the

principle of convex optimization and of the algorithms optimizing convex or non-convex functions which will be used in the following chapters.

The chapter 3 presents our first contribution: the determination of the HR-pQCT convolution kernel from couples of experimental images. The method is based on the minimization of mutual information. In order to evaluate the performance of the method, the total variation (TV) regularization is applied with the determined kernel to solve the super resolution problem. Our experimental results show that the estimated kernel has similar effects as the one obtained with basic L_2 norm minimization.

The chapter 4 discusses the application of non-convex and non-smooth optimization methods in our joint super resolution/ segmentation problem. Since the histograms of the reference micro-CT image present a bimodal distribution, we have considered a variational approach combining TV regularization with a double-well potential for our super resolution problem. We then setup three different numerical schemes to solve this non-convex non-smooth optimization problem, two being based on the Alternative Direction Methods of Multipliers (ADMM) algorithm, one being based on the accelerated proximal operator. The conclusion is that the double-well potential indeed enhances the super resolution image contrast, but does not manage to recover the details of the images.

The chapter 5 investigates the application of dictionary learning methods in the super resolution problem. Our method is based on the semi-coupled dictionary learning approach in which two dictionaries are learned corresponding to the high and low resolution. Some improvements are proposed to deal with large 3D bone images. Numerical results show that the obtained image outperforms the previous methods. Some criteria (Dice or BV/TV) are significantly improved, however, the restoration of the connectivity density of bone structure is still a big challenge.

The chapter 6 reviews the deep learning super resolution techniques and presents some applications of deep learning super resolution in medical image processing. The architecture evolves from simple shadow structure to deep densely connected network. The deep learning network has an amazing ability to approximate complex functions. How to adapt the deep learning network for medical image processing is still an active research field.

The chapter 7 presents preliminary results obtained with a proposed deep learning method. In the dictionary method, the coefficients of the high and low resolution image representations on the two dictionaries are linearly linked via two matrices. Deep learning methods permits to discover more complex mapping functions. We have applied a variant of accelerated super-resolution convolutional neural network approaches with the Nvidia GPU. The proposed network achieves a visible improvement of image quality. Moreover, we attempts to solve a joint super resolution/segmentation problem with an oversampling factor 3.42, which means improving image resolution from voxel size $82\mu\text{m}$ to $24\mu\text{m}$. Thanks to NVIDIA GPU, the process for super resolution reconstruction is very fast and allows us to apply this method on entire trabecular samples.

In conclusion, we have addressed a joint super resolution/segmentation problem applied to experimental medical bone CT images. We have first considered variational methods to estimate the unknown kernel of the problem and a nonconvex and non-smooth approach to improve the results obtained with TV regularization. Then we have considered dictionary approaches as well as deep learning approaches, which outperforms the previous methods and are thus promising for the future.

Notation

General notations

x	A scalar
$\boldsymbol{x}$	A vector or matrix or tensor
$\mathcal{F}$	A function, input is scalar, when input is not scalar, it's an element-wise operation
$\mathcal{F}$	An operator or a function, input at least includes a tensor
$\mathbb{S}$	A set
$\mathbb{R}$	The set of real numbers
$\nabla \boldsymbol{x}$	The Derivative of tensor $\boldsymbol{x}$
$\Delta \boldsymbol{x}$	The Laplacian operator of tensor $\boldsymbol{x}$
$\boldsymbol{x}^T$	The transpose of $\boldsymbol{x}$
$\boldsymbol{x} \odot \boldsymbol{y}$	The element-wise product of tensor $\boldsymbol{x}$ and $\boldsymbol{y}$
$\boldsymbol{x} * \boldsymbol{y}$	The convolution of tensor $\boldsymbol{x}$ and $\boldsymbol{y}$

Convention notations

$\boldsymbol{f}^*$	The Ground truth, high resolution image
$\boldsymbol{f}$	The super resolution image
$\boldsymbol{g}$	The observed image after degradation (low resolution image)
$\boldsymbol{n}$	Additive noise
$prox$	Proximal operator
$\mathcal{A}$	An operator includes down sampling and blur effects
$\mathcal{D}_i \boldsymbol{x}$	The differential operation at voxel i of tensor $\boldsymbol{x}$
Φ^f	In dictionary learning, it denotes the dictionary of high resolution images
Φ^g	In dictionary learning, it denotes the dictionary of low resolution images
$\mathcal{X}$	A Hilbert or Banach space
$\mathcal{Y}$	A Hilbert or Banach space

List of Figures

- 1 Histogramme typique haute résolution (taille de voxel $41\mu\text{m}$) et basse résolution (taille de voxel $82\mu\text{m}$) d'images CT. Suivant (a), s_1 et s_2 correspondent aux positions des deux pics dans l'histogramme de l'image haute résolution ($s_1 = 3750$, $s_2 = 11075$). xii
- 2 Tranches 2D d'un morceau test: HR, LR et SR désignent les images haute résolution, basse résolution et super résolution. SR-TV est l'image super résolution restaurée par TV. (d)-(f) représentent des images reconstruites par LCNNM, APG et CH+TV initialisées l'image de super résolution obtenue avec TV, SR-TV. xiii
- 3 La partie de gauche est le schéma de la méthode 2D SCDL pour la reconstruction d'image de super résolution 3D; la partie de droite est le schéma de la méthode 2.5D SCDL pour la reconstruction d'image de super résolution 3D. (1.)LR est l'image à faible résolution; (2.)init SR est la super résolution initiale en 3D; (3.)(3.1)(3.1)(3.2)(3.3.3)(3.3) sont des tranches 2D coupées à partir du volume 3D dans des directions différentes (Notez que les volumes (3.) et (3.2) sont coupés horizontalement); (4.)(4.1)(4.2)(4.3)(4.3) sont des images de super résolution restaurées avec SCDL en dimension 2D; (5.) à gauche est représenté un volume empilé avec des tranches obtenues précédemment (4.); (5.) à droite prend la moyenne de (4.1)(4.2)(4.2)(4.3). . . xvi
- 4 Illustration 3D d'images limitées à une pile de 10 tranches, sur la première ligne, de gauche à droite, haute résolution, basse résolution, super résolution obtenue avec le TV; sur la deuxième ligne, respectivement CH+TV, 2D SCDL et 2.5D SCDL. L'image affichée est un crop qui vient du set d'apprentissage et est coupé dans l'échantillon A_{tibia} , les tailles des patchs 2D et 2.5D SCDL sont 12^2 xvii
- 5 Illustration de l'architecture FSRCNN-Residual network. La première partie FSRCNN(Dong et al., 2016) traite le problème super résolution, la dernière partie Residual(Kim et al., 2016a) sert à la segmentation. Les symboles $\hat{f}$ et $\hat{f}_s$ représentent les images super résolution et la sortie du réseau, respectivement. Les cubes bleus représente les features, les cubes oranges représentent les filtres. Les numéros au dessus des cubes bleus sont les nombres des filtres (d_1, s, d_2), d_2 est le nombre de filtre dans le réseau Residual sub-network, celui qui n'est pas représenté dans cette figure. Les numéros à coté des cubes oranges sont les longueurs du bord des filtres 3D. Par exemple, quand les longueurs du bord égale à 5, la taille de ce filtre est $5 \times 5 \times 5$ xxi

6	Illustration des images 2D coupées de volumes 3D en niveau de gris. Le volume vient de l'ensemble d'apprentissage. Sur la première ligne, de gauche à droite, haute résolution, basse résolution, super résolution, image TV super résolution ; sur la deuxième ligne, les images obtenues avec CH+TV, 2.5D SCDL et l'approche proposée, respectivement.	xxii
7	Illustration de tranches 2D d'échantillons entiers en 3D à la taille voxel $24\mu\text{m}$. L'échantillon testé est un tibia dans l'ensemble de test. Cette figure est utilisée pour montrer la généralisation du réseau. Les images de la colonne de gauche sont des images de niveau de gris, de haut en bas : basse résolution, haute résolution, image super résolution reconstruite par notre méthode proposée. Toutes les images de la colonne de droite sont des images binaires. Ils sont segmentés avec Otsu, de haut en bas : basse résolution, haute résolution et image super résolution reconstruite par notre méthode.	xxiii
1.1	A–D (A, B) Cross-sectional HR-pQCT images through the distal radius show two individuals with identical BMD by DXA at the ultradistal radius but substantial differences in trabecular and cortical structure. (C, D) Three-dimensional renderings of the cortical and trabecular bone regions. The intracortical porosity (highlighted in red) was segmented using software described by Burghardt et al. (Burghardt et al., 2010). HR-pQCT = high-resolution peripheral quantitative CT; BMD = bone mineral density; DXA = dual-energy X-ray absorptiometry. (Figure from (Burghardt et al., 2011))	2
1.2	Illustration of healthy and osteoporotic trabeculae. (Image downloaded from http://www.drwolgin.com/Pages/osteoporosis.aspx)	3
1.3	Trabecular thickness illustration	4
1.4	(a)3D tomographic reconstruction of a human vertebra sample. (b)2D slice extracted from a 3D reconstruction at the level indicated by a dotted line. This example shows clearly the misleading measurement of connectivity parameters on 2D images of bone structure. What was interpreted as a free end on the slice does not correspond to a real disruption of the trabecular network, but to a small perforation.(Salomé et al., 1999)	5
1.5	Coordinate for 2D parallel X-ray reconstruction. Image collected from (Gensheng, 2010).	6
1.6	Back projection illustration of a simple example	6
1.7	Fourier slice theorem for 2D image reconstruction. This image is collected from (Gensheng, 2010)	7
1.8	Detector brings a line in Fourier domain. Once these lines covered all the Fourier space, the original image is reconstructed.	7
1.9	HR-pQCT imaging illustration	8
1.10	Representative images of the trabecular segmentation for both μ -CT and HR-pQCT images. A) Gray scale images acquired by the respective scanners. B) Trabecular segmentation of a single slice. C) 3D representation of the trabecular volume. http://radiology.ucsf.edu/research/labs/musculoskeletal-bquantitative-bimaging/labs/kazakia	9

1.11	Trabecular thickness (A) and separation (B).(Gashti et al., 2012)	11
1.12	a) Results of local topological analysis (LTA) done on the whole trabecular region of a control patient's distal radius. The following color code is used for illustration: pink: pixels of rod-like trabeculae; green: pixels of plate-like trabeculae. b) A portion of the trabecular bone was extracted to better display the plate-like and rod-like trabeculae. Image from (Peyrin et al., 2010)	13
2.1	Illustration of Tikhonov, Lasso or Elastic net regularization. The image is drawn by Joseph E. Gonzalez. Website: http://www.ds100.org/sp17/assets/notebooks/linear_regression/Regularization.html	23
2.2	Visual comparison of factor of 2 super-resolution results. The upper row shows the SR results of the image Lena. The lower row shows the SR results of the image House. The SCDL and the beta process are two approaches proposed in (Wang et al., 2012) and (He et al., 2013), respectively. The figure is adapted from (He et al., 2013)	28
2.3	Multi Layer Perceptron (MLP) neuron network. In the training stage, input and output are fed with training data, the network is optimized by minimizing a user-selected function with respect to parameters within the network. Each solid circle represents a neuron which consists of linear transformation and nonlinear activation function.	35
2.4	Illustration of three activation functions: ReLU, PReLU and sigmoid activation functions. The coefficient for the negative parts in ReLU is zero, whereas in PReLU, it is parametrized by a . Sigmoid function is used for two label classification.	36
3.1	Demonstration of the undersampling operator $\mathcal{U}$ with factor equal to 2.	40
3.2	illustration of binarized high resolution, low resolution, L2 kernel restored image and MI kernel restored image.	46
3.3	Evolution of PSNR as a function of the number of iteration for the two kernels.	46
3.4	Evolution of DICE as a function of the number of iteration for the two kernels.	47
3.5	Evolution of BV/TV as a function of the number of iteration for the two kernels.	47
3.6	Evolution of connectivity density as a function of the number of iteration for the two kernels.	48
3.7	Evolution of data fitting term as a function of iteration. In order to show the whole evolution of the data fitting term, this figure didn't stop iteration by relative changes, which is $\ \mathbf{f}^{n+1} - \mathbf{f}^n\ _2 / \ \mathbf{f}^{n+1}\ _2 \leq 10^{-3}$.	48
4.1	Typical histogram of high resolution (voxel size $41\mu\text{m}$) and low resolution (voxel size $82\mu\text{m}$) CT images. According to (a), s_1 and s_2 corresponds to the location of two peaks in the histogram of the high resolution image ($s_1 = 3750$, $s_2 = 11075$).	50
4.2	doubled-well potential function: $\mathcal{W}(\mathbf{f}_i) = (\mathbf{f}_i - 0.5)^2(\mathbf{f}_i - 2.5)^2$	51

4.3	2D slices from the first crop (A): HR and LR denote high resolution image and low resolution image respectively. SR-TV is the super resolution image restored by TV. SR-TV-LCNNM, SR-TV-APG and SR-TV-CH+TV represent images reconstructed by LCNNM, APG and CH+TV initialized with SR-TV.	57
4.4	For the first crop (A): (a)-(e) illustrate the evolutions of PSNR, DICE, BV/TV, Conn.D and $\ \mathbf{A}\mathbf{f}_p - \mathbf{g}\ _2^2$ as a function of the iteration number. For BV/TV and Conn.D, the high resolution value is also displayed as a dashed straight line on the graphs.	58
4.5	2D slices from the second crop (B): HR and LR denote high resolution image and low resolution image respectively. SR-TV is the super resolution image restored by TV. SR-TV-LCNNM, SR-TV-APG and SR-TV-CH+TV represent images reconstructed by LCNNM, APG and CH+TV initialized with SR-TV.	59
4.6	For the second crop (B): (a)-(e) illustrate the evolutions of PSNR, DICE, BV/TV, Conn.D and $\ \mathbf{A}\mathbf{f}_p - \mathbf{g}\ _2^2$ as a function of the iteration number. For BV/TV and Conn.D, the high resolution value is also displayed as a dashed straight line on the graphs.	60
5.1	The part on the left is the flowchart of 2D SCDL for 3D volume reconstruction; the part on the right is the flowchart of 2.5D SCDL for 3D volume reconstruction. (1.)LR is the low resolution image; (2.)init SR is the initial super resolution in 3D; (3.)(3.1)(3.2)(3.3) are 2D slices cut from 3D volume in different directions (Note that (3.) and (3.2) volumes are cut horizontally); (4.)(4.1)(4.2)(4.3) are super resolution images restored with SCDL in 2D dimension; (5.) on the left is a volume stacked with slices obtained previous (4.); (5.) on the right takes the average of (4.1)(4.2)(4.3).	66
5.2	2D slices cut from 3D samples.	67
5.3	Illustration of atoms in sorted dictionaries, patch size is 10^2 , atom number equals to 400. Atoms are arranged in the decreasing order of $\mathbf{M}$	69
5.4	The plot of $\mathbf{M}$ (dashed line, blue) and of its forward derivative (plain, red) as a function of the sorted dictionary.	70
5.5	Evolution of two criteria, PSNR (dashed, blue) and DICE index (plain, red) when using different numbers of atom pairs.	70
5.6	3D rendering of images limited to a stack of 10 slices, on the first row, from left to right, high resolution, low resolution, TV super resolution; on the second row, CH+TV, 2D SCDL and 2.5D SCDL, respectively. The displayed image is a training crop cut from the sample A_{tibia} , the patch sizes of 2D and 2.5D SCDL are both 12^2	72
5.7	3D rendering of images limited to a stack of 10 slices, from left to right, on the first row, high resolution, low resolution, TV super resolution, on the second row, CH+TV, 2D SCDL and 2.5D SCDL, respectively. Images is a test crop cut from the sample H_{radius} . The patch sizes of 2D and 2.5D SCDL are both 12^2	73
6.1	Illustration of SRCNN architecture, this figure is collected from (Dong et al., 2014)	78

6.2	Figure adapted from (Tai et al., 2017). (a) VDSR (Kim et al., 2016a). The skip connection between the input and output ensures the stability of the network, particularly for deep network. (b) DRCN (Kim et al., 2016b). The dashed green block is a recursive block. The convolution layers within the recursive block share the same parameters. (c) DRRN (B=1, U=2) (Tai et al., 2017). B denotes the number of recursive blocks and U is the number of recursive units in a recursive block. The green dashed block denotes a residual unit, the red dashed block represents a recursive block. Inside of a recursive block, the convolution layers in the same color share the same parameters. (d) DRRN (B=6, U=3) (Tai et al., 2017). The orange 'RB' block denotes a recursive block.	78
6.3	Illustration of ESPCN architecture. This figure is collected from (Shi et al., 2016)	80
6.4	Illustration of FSRCNN architecture. This figure is collected from (Dong et al., 2016).	80
6.5	Illustration of 8 images in the set5 as well as set14. The images on the top row are 4 images from the Set5. The images on the bottom row are 4 images from the Set14.	81
6.6	Illustration of EDSR network, the plot is collected from (Lim et al., 2017). The ResBlock is stacked by a convolution layer, Relu activation, convolution layer and a constant scaling layer. The constant scaling layer simply rescales the residual with the purpose of keeping the stability of the network when the number of filters exceeded 1000 (Szegedy et al., 2017).	82
6.7	Illustration of LapSR, this plot is collected from (Lai et al., 2017). Red arrows refers to convolution layers, blue arrows are transposed convolutions, green arrows are element-wise addition.	83
6.8	Illustration of SRGAN, the diagram is collected from (Ledig et al., 2016). k represents the size of filter, n as the number of feature maps, s as the stride for each convolutional layer.	83
6.9	Illustration of densely connected CNN, this figure is collected from (Huang et al., 2017)(denseNet).	84
6.10	Illustration of the architecture of SRdenseNet, this diagram is collected from (Tong et al., 2017). (a) uses dense block to enhance image features then feed the enhanced features into deconvolution layers for upsampling. (b) long skip connection preserves low level features and enables the dense blocks to learn high level features. (c)The bottleneck combines both low and high level features and feed the results into deconvolution layers. . . .	85
6.11	Comparison of different network blocks, this image is collected from (Zhang et al., 2018). (a) Residual block in (Lim et al., 2017). (b) Dense block in SRDenseNet (Tong et al., 2017). (c) Residual dense block (Zhang et al., 2018)	86
6.12	The illustration of the residual dense network, the plot is collected from (Zhang et al., 2018).	86
6.13	The proposed framework for MRI super resolution, this figure is collected from (Zhao et al., 2018).	88
6.14	Framework of 3D DCSRN, this plot is collected from (Chen et al., 2018). .	89

6.15	The proposed T-L network segmentation and super resolution tasks, this figure is collected from (Oktay et al., 2018).	89
7.1	Illustration of the proposed FSRCNN-Residual network. The former sub-network FSRCNN(Dong et al., 2016) tackles the super resolution task, the sequential sub-network Residual(Kim et al., 2016a) serves for segmentation purpose. The symbols $\mathbf{f}$ and $\mathbf{f}_s$ denote super resolution images and output images, respectively. The blue cubes denote features, orange cubes represent filters. The numbers above the blue cubes are the number of filters (d_1, s, d_2), d_2 is the number of filters in the Residual sub-network which is not depicted on this figure. The numbers below orange cubes are the length of 3D filter edges. For instance, when the length of filter edge is 5, the size of filter is $5 \times 5 \times 5$	92
7.2	Illustration of a 2D slice of 3D image of the training set in gray level, the upsampling factor is 2. The label of the sample is 'C0008233-D0000812'. On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed FSRCNN-Residual, respectively.	95
7.3	Illustration of a 2D slice of 3D image of the training set in binary, the upsampling factor is 2. The label of the sample is "C0008233-D0000812". On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.	96
7.4	Illustration of a 2D slice of 3D image of the test set in gray level, the upsampling factor is 2. The label of the sample is 'C0008231-F0000814'. On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.	97
7.5	Illustration of a 2D slice of 3D image of the test set in binary, the upsampling factor is 2. The label of the sample is "C0008231-F0000814". On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.	98
7.6	Illustration of the proposed FSRCNN-Residual network. The former sub-network FSRCNN(Dong et al., 2016) tackles the super resolution task, the sequential sub-network Residual(Kim et al., 2016a) serves for segmentation purpose. The symbols $\mathbf{f}$ and $\mathbf{f}_s$ denote super resolution images and output images, respectively. The blue cubes denote features, orange cubes represent filters. The numbers above the blue cubes are the number of filters (d_1, s, d_2), d_2 is the number of filters in the Residual sub-network which is not depicted on this figure. The numbers below orange cubes are the length of 3D filter edges. For instance, when the length of filter edge is 5, the size of filter is $5 \times 5 \times 5$. It's noteworthy that the size of the deconvolution filter is $8 \times 8 \times 8$, stride of convolution is 3, no padding involved.	99
7.7	Demonstration of convolution with stride: a 4×4 signal is convolved with a 2×2 filter with stride 2.	100

7.8	Illustration of a 2D slice of 3D volume at voxel size $24\mu\text{m}$. The crop is cut from test set. The label of sample is "C0008230-G0000818". The images on the top row are gray level images, from left to right: low resolution, high resolution, our method; All the images on the bottom row are binary images segmented with Otsu, from left to right: low resolution, high resolution and our proposed method.	101
7.9	Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The sample is a tibia in the training set. The label of the sample is "C0008233-D0000812". This figure is used to show the best performance of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.	103
7.10	Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The sample is a tibia in the test set. The label of the sample 'C0008233-G0000812'. This figure is used to show the generalization property of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.	104
7.11	Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The tested sample is a radius in the test set. The label of the image is 'C0008231-F0000814'. This figure is used to show the generalization property of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu threshold, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.	105
7.12	Illustration of the effects of mask on a tibia sample. The label of the sample is "C0008233-D0000812". The images on the left column include cortical as well as trabecular region. The images on right column only displays the trabecular structures abstracted by the mask. From top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method.	107
7.13	Illustration of the effects of mask on a radius sample. The label of the sample is "C0008230-G0000818". The images on the left column include cortical and trabecular region. The images on the right column displays the trabecular structures abstracted by the mask. From top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method.	108
7.14	Summary of BV/TV over the entire dataset. The first 4 samples are in the training set.	109

LIST OF FIGURES

7.15	Illustration of samples with inferior BV/TV estimation. The trabecular samples on the left column is 'C0008230-D0000818', on the right is 'C0008233-F0000812'. From up to bottom, the images are binary low resolution, high resolution and super resolution images.	111
7.16	Illustration of samples with inferior BV/TV estimation. The trabecular samples on the left column is 'C0008231-G0000814', on the right is 'C0008231-F0000814'. From up to bottom, the images are binary low resolution, high resolution and super resolution images.	112
7.17	The histogram distribution of the sample 'C0008231-F0000814' at low resolution scale. It's noteworthy that this histogram display a bimodal shape since the cortical region is taken account of. There will be only one peak observed in the histogram if no cortical region is considered.	113
7.18	Comparison of super resolution images obtained with different scale factor in 'Normalization'. All the images are binary.	114
8.1	Diagram of registration	124

List of Tables

5.1	Image super-resolution criteria for the training crop cut from A_{tibia} obtained with different patch sizes. The over completeness factor is 4. 2.5D SCDL reconstructs 3D images with 2D dictionaries	71
5.2	Image super-resolution criteria for the test crop cut from D_{radius} obtained with different patch sizes. The over completeness factor is 4. 2.5D SCDL reconstructs 3D images with 2D dictionaries	71
5.3	Criteria for 5 test crops obtained with 2.5 SCDL.	71
5.4	criteria of the training crop cut from the sample A_{tibia} obtained with different super resolution methods. 2.5D SCDL reconstructs 3D images with 2D dictionaries	73
5.5	criteria of the test crop cut from the sample H_{radius} obtained with different super resolution methods. 2.5D SCDL reconstructs 3D images with 2D dictionaries	74
6.1	Comparison of PSNR different methods for natural image sets set5 and set14, the low resolution images are generated with BI degradation model. The upsampling factor is 4. The sign '+' represents that the approach has been boosted with self-assemble. The results are collected from correspondent publications.	87
7.1	Comparison of different super resolution approaches. The volume is a $200 \times 200 \times 200$ crop cut from the sample 'C0008233-D0000812' in the training set, the upsampling factor is 2.	95
7.2	Comparison of different super resolution approaches. The tested volume is a $200 \times 200 \times 200$ crop cut from the sample 'F0000814-C0008231' in the test set.	97
7.3	Comparison of different super resolution approaches. The evaluated subject is a $410 \times 410 \times 123$ crop cut from the sample 'C0008230-G0000818' in the test set.	101
7.4	Summary of Dice and BV/TV over the entire dataset. The first 4 samples are in the training set, the other 10 are in the test set. The cells in yellow highlights the poor estimation of BV/TV on super resolution samples. . . .	110
7.5	Comparison of super resolution images obtained with different scale factors in the 'normalization'. The investigated sample is 'C0008233-F0000812'. . .	115

7.6	Analysis of super resolution BV/TV error with the histogram of low resolution input images. The 'peak 1' and 'peak 2' are the two most obvious peaks positions (intensity) in the histogram of low resolution images. When the cortical regions are thin or non-homogeneous, the histograms do not have bimodal shape, in this case, the correspondent cells in the 'peak 2' column are marked with 'N', similarly for the column 'max/peak 2'. The column 'max' records the max intensity value on the low resolution image. 'max/peak 2' reflects the relative position between the second peak in the histogram and the max value in the low resolution images. 'SR BV/TV error' is the relative error of the BV/TV of super resolution images with respect to the ground truth. When its sign is '+', it means that the BV/TV of the super resolution image is overestimated, the sign '-' means that the BV/TV is underestimated. Cells marked in orange highlight relevant information of samples whose super resolution BV/TV are highly underestimated.	115
8.1	Different resolution images are involved in a sample registration process. Here we analyse sample 'C0008230-D0000818' and the table summarized topological criteria of different resolution images. This table only considered Bone volume and Connectivity, but not BV/TV or Conn.D, since BV/TV or Conn.D are computed based on total volume of the sample. For a fixed sample at different resolution, firstly, they share the sample total volume, and secondly, if we generate mask for different resolution images to obtain total volume corresponding to different resolution, errors may be introduced.	126
8.2	Parameter synthesis after new registration	129

Contents

Résumé étendu	i
Notation	xxix
List of figures	xxxviii
List of tables	xl
1 Context and Motivations	1
1.1 Osteoporosis	1
1.2 Bone micro architecture	3
1.3 Medical imaging for bone micro-architecture	4
1.3.1 CT image reconstruction	5
1.3.2 HR-pQCT	8
1.3.3 Micro-CT	9
1.3.4 Synchrotron Radiation Micro CT	10
1.4 Bone parameters	11
1.5 Pros and cons of HR-pQCT images	13
1.6 Summary	15
2 Introduction to super resolution methods and mathematical foundations	17
2.1 The inverse problem model for super resolution	17
2.2 Different approaches for super resolution	18
2.2.1 Interpolation in spatial domain	18
2.2.2 Super resolution in frequency-wavelet domain	20
2.2.3 Regularization methods for super resolution	21
2.2.4 Dictionary learning	25
2.3 Convex optimization	27
2.3.1 Proximal operator	28
2.3.2 Alternate direction method of multiplier(ADMM)	30
2.4 Nonconvex optimization methods	31
2.4.1 Accelerated proximal gradient method	32
2.4.2 Ipiano: inertial proximal approach	32
2.4.3 Linearly constrained nonconvex nonsmooth minimization	32
2.4.4 DC programming	33
2.5 Deep learning introduction	34

2.5.1	Generalities	34
2.5.2	Network units	34
2.5.3	CNN	36
2.5.4	Optimization of the network	36
2.6	Discussion	37
3	Determination of the kernel	39
3.1	Introduction	39
3.2	The inverse problem of the blurring kernel estimation	40
3.2.1	L_2 norm minimization	41
3.2.2	Mutual Information based approach	41
3.3	Super resolution with the TV regularization	43
3.4	Numerical experiments	44
3.4.1	Simulation details	44
3.4.2	Criteria for image quality	45
3.4.3	Numerical results	45
3.5	Conclusion	47
4	Nonconvex and nonsmooth optimization to enhance the contrast of images	49
4.1	Introduction	49
4.2	The inverse problem	50
4.3	Algorithms for nonconvex and nonsmooth optimization	51
4.3.1	Linearly Constrained Nonsmooth and Nonconvex Minimization method (LCNNM)	51
4.3.2	Accelerated Proximal Gradient (APG)	53
4.3.3	A combination of Cahn-Hilliard model and Total Variation (CH+TV)	53
4.4	Application to bone CT images	56
4.4.1	Methodology	56
4.4.2	Results	58
4.5	Discussion and Conclusions	60
5	Investigation of semi-coupled dictionary learning in 3D super resolution HR-pQCT imaging	63
5.1	Introduction	63
5.2	Semi-coupled dictionary learning approach	64
5.2.1	Dictionary learning stage	64
5.2.2	Super resolution stage	65
5.2.3	Implementation details	66
5.2.4	SCDL scheme for 3D trabecular bone images	67
5.3	Results	67
5.3.1	Selection of 2D atoms	68
5.3.2	Results on 3D images	69
5.3.3	Comparison with TV based methods	71
5.4	Conclusion	73
5.5	Discussion	74

6	Deep learning for super resolution	77
6.1	The state of the art on deep learning for super resolution	77
6.1.1	Deep learning network for super resolution problems	77
6.1.2	Comparison of the performances of the different networks	87
6.2	The application of deep learning methods in medical images super resolution problems	87
7	A proposed network for our super resolution problem: FSRCNN-Residual	91
7.1	FSRCNN-Residual Architecture	92
7.1.1	The FSRCNN sub-network	92
7.1.2	The Residual sub-network	93
7.1.3	The loss functions	93
7.2	Numerical experiments	94
7.2.1	Dataset and criteria	94
7.2.2	Super resolution results for crops	94
7.2.3	Discussion	96
7.2.4	Conclusion	98
7.3	The application of FSRCNN-Residual in super resolution problem with upsampling factor 3.42	98
7.3.1	The adapted FSRCNN-Residual	99
7.3.2	Numerical results	100
7.3.3	Computation trabecular bone structure parameters	106
7.3.4	Image quantification	109
7.4	Conclusion and discussion	116
8	Conclusions and Perspectives	117
	Annexes	123

CONTENTS

Chapter 1

Context and Motivations

Our study is motivated by the development of imaging methods in the context of the diagnosis and follow up of osteoporosis. This disease is related to bone fragility and is responsible of fractures and invalidity. In this context, X-ray imaging is used at different scales to assess bone properties *ex vivo* and *in vivo*.

In this chapter, we first explain what is osteoporosis in Sect.1.1, then we briefly describe the anatomy of bone and why we need to investigate bone micro-architecture in Sect.1.2. In Sect.1.3, different techniques for bone micro-architecture imaging are introduced. And in Sect.1.4, we presents a number of bone parameters used to quantify bone quality based on the micro-architecture. The objective of this thesis is presented in Sect.1.5, including how to model the problem and what types of data we deal with.

1.1 Osteoporosis

Osteoporosis is one of the most common bone diseases people have to face today. It is characterized by a loss of bone mass and deterioration of bone micro-architecture, which increases the fragility of bone and augments the risk of fracture ([Harvey et al., 2010](#)). Early osteoporosis is asymptomatic and should be diagnosed before a fracture happens. The rate of fracture increases with advancing age due to the degradation of bone micro-architecture and the loss of bone mass. People tend to loss bone mass after 50 years old, especially women after menopause. With the development of health care and medical treatment, the increase of life expectancy may give us more time to develop a disease or to fight against osteoporosis. Osteoporosis is not a fatal disease, however, the fractures it induces may have serious complications (damage of vessels and nerves, infections, stiffness), sometimes even fatal threats ([Carpintero et al., 2014](#)). Once a fracture occurred, patients need long nursing time to recover, which is also a heavy burden in term of personal life quality and economics ([Johnell and Kanis, 2006](#)).

Three main fracture sites in human are hip, vertebrae and forearm ([on Prevention et al., 2003](#)). Hip fracture requires 30 days hospitalization in average during which the rate of morbidity and mortality increase ([Kamel et al., 2000](#)). Meanwhile, only 30% of patients recover their original life quality. And patients who have suffered from hip structures have a higher risk of suffering from another hip fracture later ([Schroder et al., 1993](#)). Vertebrae fracture rarely needs hospitalization. It has a longer impact of patient life quality rather than increasing the mortality as hip fracture. Forearm fracture is less

harmful than hip or vertebrae fractures. Yet, it may also damage the surrounding nerves and vessels or induce other complications, and takes 3-6 months to fully heal. Patients after forearm fracture may have a higher risk of another fracture later ([Silman, 1995](#)).

Many researchers have attempted to analyze osteoporosis from different aspects, such as age, gender, habits, etc. A large proportion of people after 75 years old suffer from osteoporosis, for both genders. In fact, elderly population is a population at high-risk for osteoporosis. The situation becomes even more severe in the regions where the size of elderly population increases fast ([on Prevention et al., 2003](#)). Osteoporosis may also occur to a small subset of postmenopausal women from 51 to 65 years old, the same syndrome has been observed to the men at comparable age but with less frequency ([Riggs and Melton, 1983](#)). It is explained in ([on Prevention et al., 2003](#)) that women has a lower peak bone mass than men and the hormonal changes during menopause increases the risk of osteoporosis. Nevertheless, Kauffman *et al.* reveals that the osteoporosis is underestimated for men ([Kaufman et al., 2013](#)). Even though in average, men tend to have bone fracture 5-10 years later than women ([Campion and Maricic, 2003](#)), the rate of morbidity and mortality is higher than women after hip fracture ([Khosla et al., 2008](#)).

In addition, cigarette smoking, physical exercises, diets, lack of calcium or vitamin D may also be the causes of osteoporosis([Weaver et al., 2016](#),[Hollenbach et al., 1993](#)).

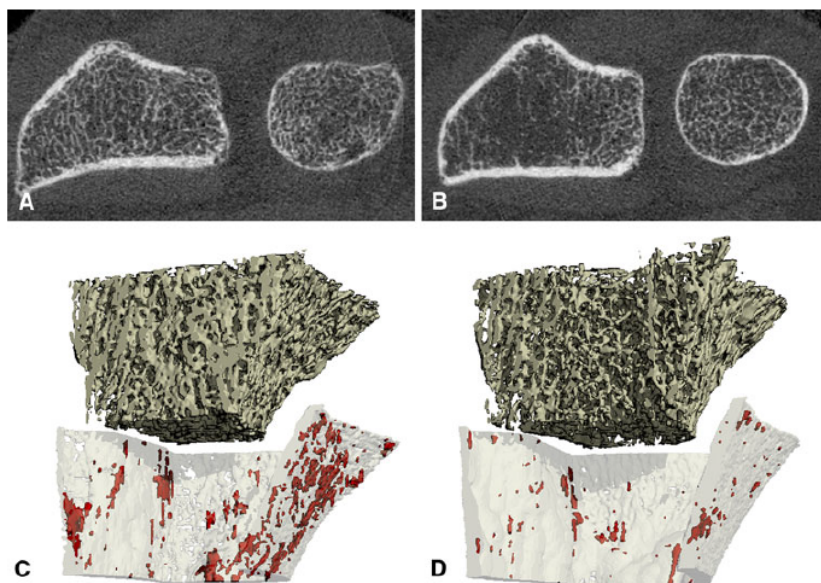


Figure 1.1: A–D (A, B) Cross-sectional HR-pQCT images through the distal radius show two individuals with identical BMD by DXA at the ultradistal radius but substantial differences in trabecular and cortical structure. (C, D) Three-dimensional renderings of the cortical and trabecular bone regions. The intracortical porosity (highlighted in red) was segmented using software described by Burghardt *et al.* ([Burghardt et al., 2010](#)). HR-pQCT = high-resolution peripheral quantitative CT; BMD = bone mineral density; DXA = dual-energy X-ray absorptiometry. (Figure from ([Burghardt et al., 2011](#)))

Generally speaking, it is bones that shape our body, support our daily activities, protect our organs and serve as a mineral storage for the blood mineral level. It is so crucial for the health, the methods to quantify the quality of bone is an important

question.

Bone mineral density (BMD) is a criterion to diagnose osteoporosis and is measured by dual-energy X-ray absorptiometry (DXA). DXA has two X-ray beams with different energy. The BMD can be estimated by eliminating the absorption of soft tissue. However, this technique can not provide an accurate *in vivo* estimation since it can only measure objects with two X-ray absorption disparate substances, which is not the case for *in vivo* imaging (Bolotin, 2007). Researches showed that BMD only represents 70-75% of bone strength variability, the remaining part is related to bone micro-architecture and tissue composition. Figure 1.1 illustrates these two types of bone micro structure images by high resolution peripheral quantitative CT (HR-pQCT): the samples with the same BMD may have different bone micro-architectures, hence distinct bone strength to resist fracture. This example shows that it is not reliable to determine bone fragility only by referring to BMD. The bone micro architecture serves as a crucial factor influencing the bone strength as well. We introduce it in the next section.

1.2 Bone micro architecture

Conventionally, bone is classified into two categories: compact (cortical) bone and trabecular (spongy or cancellous) bone. Cortical bones are very compact and dense. They form the highly compact outer layer of bone. Trabecular bones form the inner layer of bone with a lower density. Some cavities are irregularly distributed within the spongy structure of trabecular bones. These cavities contain red bone marrow, which contributes to blood generation.

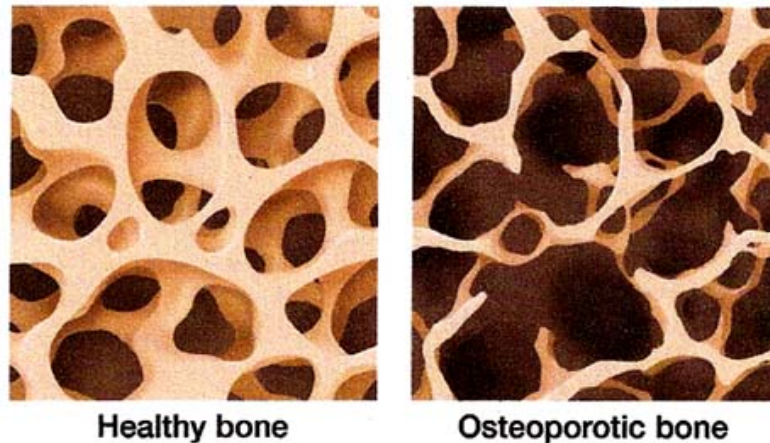


Figure 1.2: Illustration of healthy and osteoporotic trabeculae. (Image downloaded from <http://www.drwolgin.com/Pages/osteoporosis.aspx>)

Trabecular bones usually turn over with a faster speed than cortical bones (Rodan, 1992), thus it is more sensitive to the loss of bone mass. Fig.1.2 illustrates two types of trabecular bone structure: the healthy bone has a nest-like highly connected structure, while the osteoporotic bone has a lower bone mass with poor connectivity. By imaging the micro-architecture of trabeculae, many crucial bone parameters including bone mass can

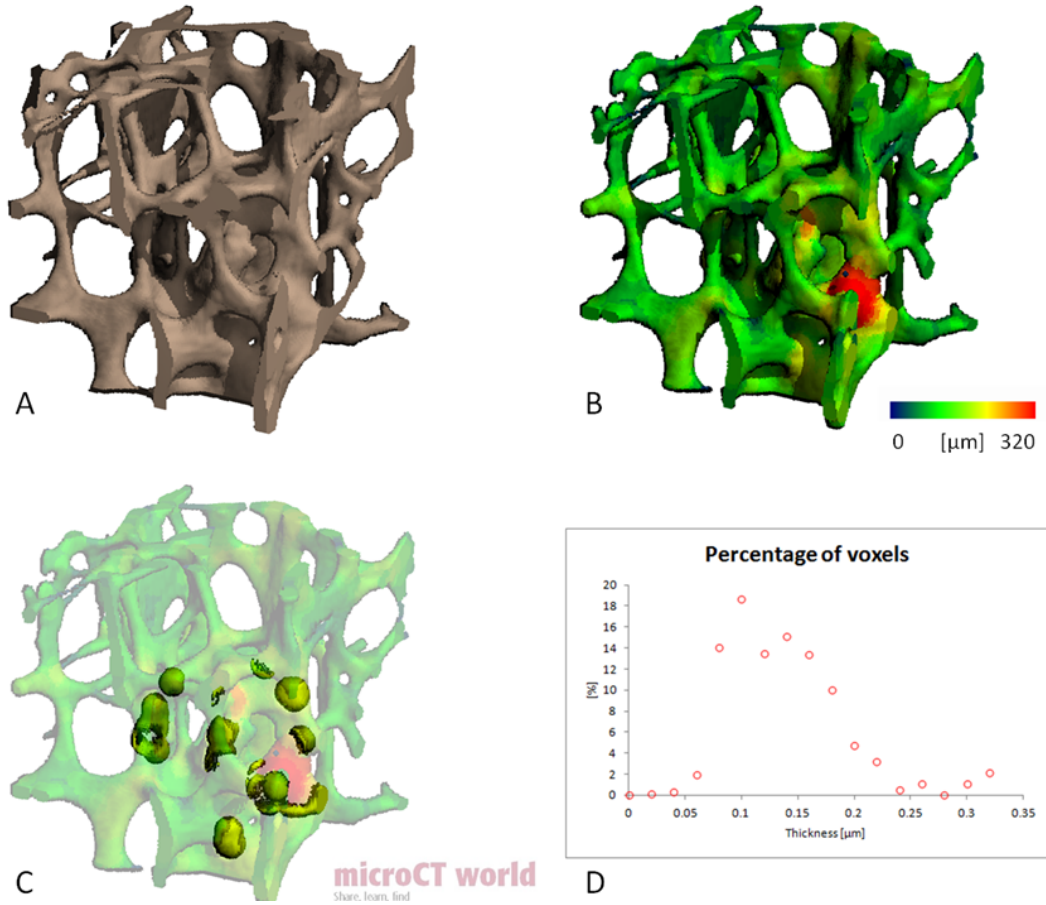


Figure 1.3: (A) Trabecular bone sample ($3 \times 3 \times 3 \text{ mm}^3$). B) Same as A), where the colors code the local thickness (0 – 320 μm). C) Color coded trabecular bone in semi-transparent and all voxels that are assigned to a value of 200 μm are shown in yellow. D) Histogram of thickness distribution. Images downloaded from the site <http://microctworld.net/trabecular-bthickness-btb-bth-btrabecular-bspacing-btb-bsp-btrabecular-bnumber-btb-bn/>

be estimated so that the quantification of bone quality becomes more accurate. Figure 1.3 presents the thickness distribution of trabecular structure. Such parameters characterize the strength of bone micro-architecture, hence facilitate and improve the diagnosis and treatment of osteoporosis or other skeleton diseases (Brandi, 2009).

3D micro-CT ($\mu\text{-CT}$) imaging has been developed in the context of bone imaging in the 1990's and has progressively replaced histology, based on 2D microscopic imaging. 3D imaging can reflect the situation of trabeculae more specifically compared to 2D. As Fig.1.4 shows, some connections of trabeculae are not visible in 2D so that it's diagnosed as a perforation. It's much easier and more accurate for 3D imaging to identify whether the trabecular network is disrupted and suffers from perforation (Salomé et al., 1999).

1.3 Medical imaging for bone micro-architecture

Bone micro-architecture can be investigated by Magnetic Resonance Imaging (MRI) or X-ray Computerized Tomography (CT). MRI generates images by observing the reaction

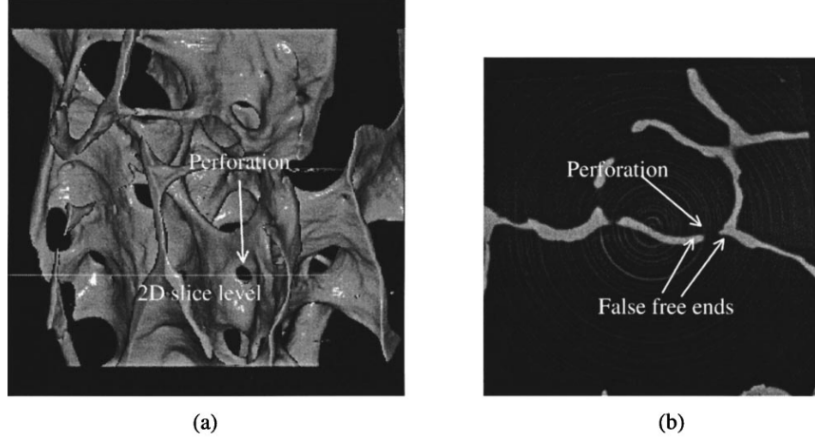


Figure 1.4: (a) 3D tomographic reconstruction of a human vertebra sample. (b) 2D slice extracted from a 3D reconstruction at the level indicated by a dotted line. This example shows clearly the misleading measurement of connectivity parameters on 2D images of bone structure. What was interpreted as a free end on the slice does not correspond to a real disruption of the trabecular network, but to a small perforation. (Salomé et al., 1999)

of protons within different tissues under a magnetic field obtained with radio waves at an appropriate resonance frequency. MRI is a non-invasive modality and it has extensive applications in clinical imaging and *in vivo* researches. However, its resolution is still limited compared with other high resolution CT techniques for trabecular bone micro architecture imaging.

1.3.1 CT image reconstruction

The development of CT originates from X-ray radiography. This technique was used first to reconstruct 2D images and later on for 3D images. The CT beam has evolved from a parallel geometry to a fan beam one then to a cone beam. For the parallel geometry, the reconstruction is based on the Radon Transform (Feldkamp et al., 1984) model.

Radon transform

Let us denote $f(x, y)$ the density in the $x - y$ plane, $\delta(\cdot)$ the Dirac function, then the Radon transform can be expressed as:

$$p(s, \theta) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x, y) \delta(x \cos \theta + y \sin \theta - s) dx dy \quad (1.1)$$

Back projection

The back projected image $b(x, y)$ can be written as

$$b(x, y) = \int_0^\pi p(s, \theta) |_{s=x \cos \theta + y \sin \theta} d\theta \quad (1.2)$$

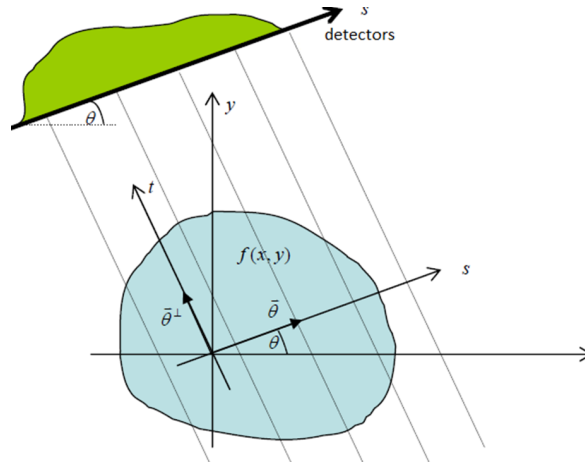


Figure 1.5: Coordinate for 2D parallel X-ray reconstruction. Image collected from (Gensheng, 2010).

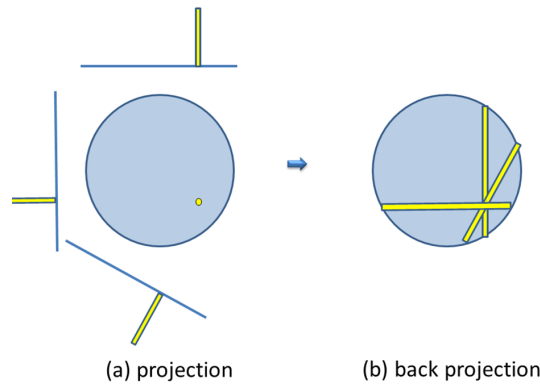


Figure 1.6: Back projection illustration of a simple example

Zen *et al.* has written a book (Gensheng, 2010) to introduce medical image reconstruction. They have an interesting metaphor to interpret this equation. $p(s, \theta)$ is the projection of object at angle θ , how to project it back? It's like spread seeds. For a fixed s and θ , we get an scalar value which is the line integration of object along the direction perpendicular to θ , but no prior for what happens within the line integration. We only see the output of the integration, but don't know how and which pixels take part in this sum operation, in another word, we don't know where to spread more seeds and where less. As a result, only one projection angle is not enough for image reconstruction, but more projections with different angles can describe images. Fig1.6 illustrates that the final image can be determined by the integration of projections at different angles, as suggested by Eq.1.2.

Fourier slice theorem

The Fourier slice theorem, also known as center slice theorem, is crucial in the theory of CT image reconstruction. It says that the 1D Fourier transform of $p(s, \theta)$ equals to the 2D Fourier transform of $f(x, y)$ at angle θ . Fig.1.7 illustrates this theory (Gensheng,

2010).

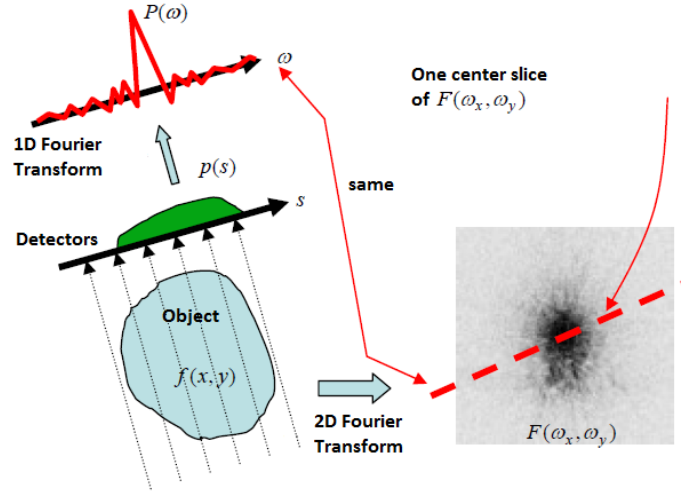


Figure 1.7: Fourier slice theorem for 2D image reconstruction. This image is collected from (Gensheng, 2010)

With this fundamental result, when a complete Fourier transform $F(\omega_x, \omega_y)$ has been estimated after 180° rotation, the reconstructed image $f(x, y)$ can be determined.

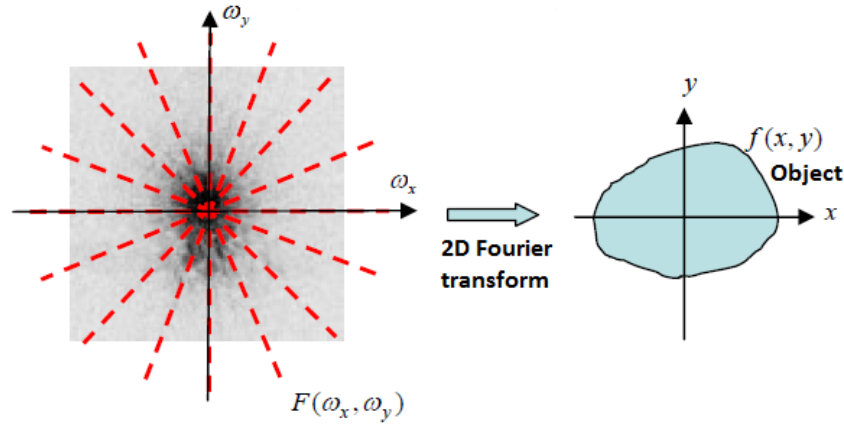


Figure 1.8: Detector brings a line in Fourier domain. Once these lines covered all the Fourier space, the original image is reconstructed.

Filtered Back projection

However, one drawback of the reconstruction method mentioned above is that image frequency spectrum has a higher density distribution near origin point than other regions. Therefore it is necessary to introduce a correction filter to reduce the over weighted effects at low frequency. Since the density of Fourier spectrum is proportional to $1/\sqrt{\omega_x^2 + \omega_y^2}$, one solution is to filter the back projection 2D image, or to filter 1D projected data before back projection. The latter one is known as Filtered Back Projection (FBP) algorithm.

Ramp filter is a basic filter in FBP. In Fourier domain it is expressed as $|\omega|$. Apart from ramp filter, many filters are used for CT reconstruction in the literature, such as Ramachandran - Lakshminarayanan filter, Shepp-Logan filter, cosine filter, etc.

With a chosen filter $h(s)$, the parallel beam reconstruction formula is described as:

$$f(x, y) = \int_0^\pi [h(s) * p(s, \theta)]|_{s=x \cos \theta + y \sin \theta} d\theta \quad (1.3)$$

1.3.2 HR-pQCT

HR-pQCT is a device designed for *in vivo* imaging of 3D bone micro-architecture at peripheral sites in human, see Fig.1.9.



Figure 1.9: HR-pQCT imaging illustration

HR-pQCT imaging is performed with XtremeCT system (Scanco Medical), with a high isotropic resolution, and a voxel size of $41\mu\text{m}$, $82\mu\text{m}$ or $123\mu\text{m}$. *In vivo* HR-pQCT imaging is generally performed at a voxel size of $82\mu\text{m}$, while *ex vivo* imaging can be done at a voxel size down to $41\mu\text{m}$. Figure 1.10 presents the HR-pQCT images at different resolutions. The new generation (XtremeCT II) for *in vivo* imaging can reach $61\mu\text{m}$ and *ex vivo* $30\mu\text{m}$. HR-pQCT, even though is not widely applied in the clinic level, are increasingly involved in clinic research and basic studies (Figueiredo et al., 2017, KK et al., 2017). It allows researchers to quantitatively analyze *in vivo* bone micro-architecture and favors the development of clinic diagnosis as well as followed-up treatment. These images are used to compute a number of bone parameters that will be described in Sect.1.4. These parameters assist researchers and clinicians to estimate the quality of bone and to quantify the changes within bone micro-architecture.

As previously mentioned, the main sites of osteoporosis include wrist, spine and hip. While HR-pQCT provides currently the best resolution to investigate bone micro-architecture *in vivo*, it is still constrained to peripheral sites for *in vivo* assessment and can't be applied for other sites like spine or hip imaging in clinical application. Moreover, the resolution of HR-pQCT is of the same order than the trabecular bone size, as highlighted by (Geusens et al., 2014). A higher resolution of HR-pQCT is expected for further researches.

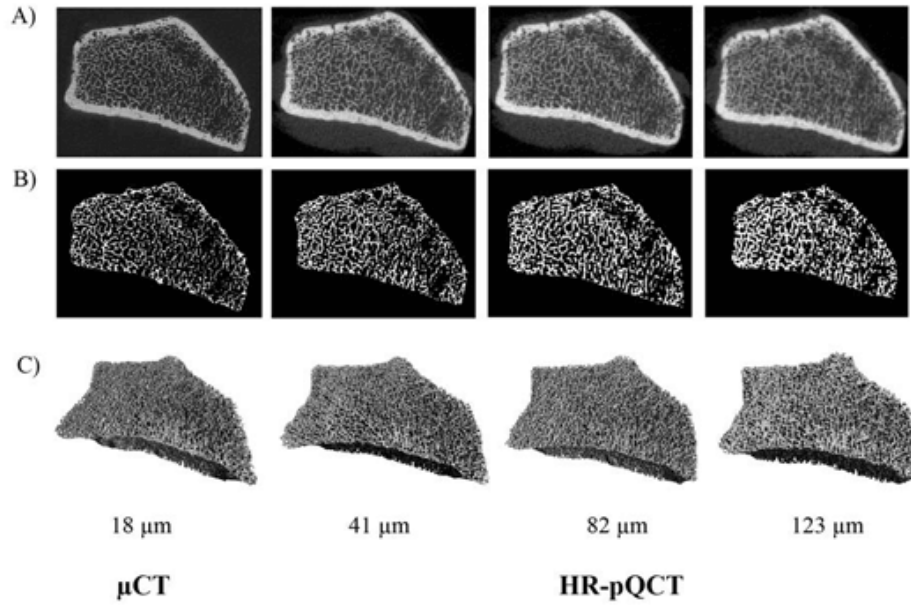


Figure 1.10: Representative images of the trabecular segmentation for both μ -CT and HR-pQCT images. A) Gray scale images acquired by the respective scanners. B) Trabecular segmentation of a single slice. C) 3D representation of the trabecular volume. <http://radiology.ucsf.edu/research/labs/musculoskeletal-bquantitative-bimaging/labs/kazakia>

1.3.3 Micro-CT

Micro-CT (μ -CT) is an imaging system which is nondestructive for bone samples and noninvasive for animal models. It uses X-ray tube and the projection is often realized with cone beam geometry. Despite being unable to scan *in vivo* human tissue, μ -CT has an excellent performance for scanning *in vivo* small animals and *ex vivo* cadaver samples as well as biopsies. It provides images with spatial resolution up to the micrometer level (Peyrin and Engelke, 2012). The higher the resolution is, the smaller the field of view (FOV) is. For example, *in vivo* small animal scanning needs a large FOV, as a consequence, the image spatial resolution varies from $40\mu\text{m}$ to $100\mu\text{m}$, which is relatively large in μ -CT resolution range. In this type of scanning, it is the gantry that rotates. Contrary to *in vivo* scans, the *ex vivo* images for cadaver samples and biopsies have a high spatial resolution, and therefore the FOV is smaller. In cone beam systems, μ -CT works with a fixed gantry and allows to acquire images at different spatial resolutions by adapting the source detector distance.

Fig.1.10 illustrates μ -CT and HR-pQCT images at different spatial resolution. Since μ -CT can provide accurate parameters in topological and morphological analysis, it is usually regarded as a gold standard for HR-pQCT. Researchers tend to study HR-pQCT image quality by comparing bone parameters extracted from HR-pQCT and μ -CT images of the same sample. For instance, research work in (Liu et al., 2010) and (MacNeil and Boyd, 2007) assessed micro-architecture of human distal bone samples via μ -CT and HR-pQCT devices. They aimed to find the correlation of micro-structure measurements

Table 1, HR-pQCT parameters with units (Odgaard, 1997)

	Abbreviation	Description	Standard unit
Metric measures			
Total volume	TV	Volume of entire region of interest	mm ³
Bone volume	BV	Volume of region segmented as bone	mm ³
Bone surface	BS	Surface area of the region segmented as bone	mm ²
Bone volume ratio*	BV/TV	Ratio of bone volume to total volume in region of interest	%
Bone surface ratio	BS/BV	Ratio of bone surface area to bone volume	%
Trabecular thickness*	Tb.Th	Mean thickness of trabeculae	mm
Trabecular thickness SD	Tb.Th.SD	Measure of homogeneity of trabecular thickness	mm
Trabecular separation*	Tb.Sp	Mean space between trabeculae	mm
Trabecular separation SD	Tb.Sp.SD	Measure of homogeneity of trabecular spacing	mm
Trabecular number*	Tb.N	Mean number of trabeculae per unit length	per mm
Cortical thickness (original)*	Ct.Th	Average cortical thickness	mm
Cortical porosity	Ct.Po	Cortical porosity	%
Bone mineral density (D100)*	BMD	Total volumetric density	mg HA/cm ³
Cortical bone mineral density (Dcomp)*	Ct.BMD	Cortical volumetric density	mg HA/cm ³
Trabecular bone mineral density (Dtrab)*	Tb.BMD	Trabecular volumetric density	mg HA/cm ³
Total bone area*	Tt.Ar	Cross-sectional area	mm ²
Non-Metric Measures			
Structural model Index	SMI	Measure of trabecular structure (0 for plates and 3 for rods)	
Degree of anisotropy	DA	1 is isotropic, >1 is anisotropic by definition; DA = length of longest divided by shortest mean intercept length vector	
Connectivity density	Conn.D	Extent of trabecular connectivity normalized by TV	mm ⁻³
Cross-sectional moment of inertia	Imin, I max	minimum and maximum moments of inertia	mm ⁴

*Standard HR-pQCT measures are indicated

obtained with these two imaging systems in order to validate the accuracy of *in vivo* HR-pQCT. The former study was performed at a distal tibia site and the latter at a distal radius site. They both conclude that most measurements of HR-pQCT and μ -CT are highly correlated, but the observation on the coherence of some individual measurements is opposite. According to Liu *et al.* in (Liu et al., 2010), they explained these discrepancies with two reasons: first, the sites of the tested samples are distinct; second, the segmentation methods are very crucial for the final analysis, which may also induce uncertainty on the final conclusion.

1.3.4 Synchrotron Radiation Micro CT

An even better resolution is obtained with Synchrotron Radiation Micro CT (SR- μ -CT). SR- μ -CT uses synchrotron as X-ray sources, with a brilliance several order of magnitude higher than the one of X-ray tubes. The projections are measured with a quasi-monochromatic beam with a high flux. Such quasi-monochromatic beams with high flux do not only avoid beam hardening effect, but also obviously shorten the time of acquisition. For example, for the same image quality with spatial resolution at micrometer level, a standard μ -CT takes 10 hours, while SR- μ -CT only uses 10 minutes (Peyrin and Engelke, 2012). During the imaging process, the gantry is fixed and the sample is rotated horizontally to cover 180 degree.

SR- μ -CT has numerous applications for the study of micro-structures of the order of μm or at nano scale (Peyrin and Engelke, 2012). SR- μ -CT can measure the mineralization of bone tissue as well. Peyrin *et al.* started to develop SR- μ -CT on beam line ID19 at the European Synchrotron Radiation Facility (ESRF) (Salomé *et al.*, 1999) to quantify trabecular bone micro architecture. Since the sample is positioned at 145m away from source, the beam can be regarded as a parallel beam. This setup has been used in many studies to image bone in small animal models (Martín-Badosa *et al.*, 2003), in human bone biopsies (Nuzzo *et al.*, 2002) or in biomaterials for instance (Peyrin *et al.*, 2007).

According to the official site of ESRF, ESRF has a various applications in nowadays research, including medicine, chemistry, culture, environment and so on. It allows to quantify and qualify tiny structures, to observe fast changes such as biological proteins.

1.4 Bone parameters

Bone parameters definitions

A number of parameters are conventionally used to characterize the bone micro-architecture. These bone parameters can be computed from HR-pQCT images and are listed in Table 1.

BV/TV reflects the ratio of bone volume (BV) to total volume(TV).

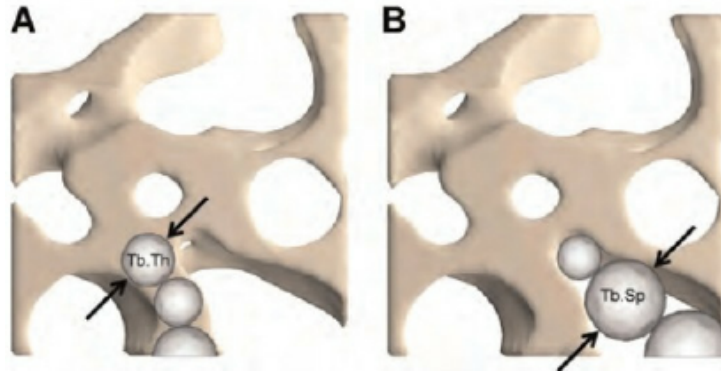


Figure 1.11: Trabecular thickness (A) and separation (B). (Gashti *et al.*, 2012)

Figure 1.11 presents the notion of trabecular thickness (Tb.Th) and separation (Tb.Sp). Tb.Th is the diameter of the maximal sphere within the structure. The Tb.Th measures the thickness of bone structure, while the Tb.Sp measures the thickness of marrow. As for Trabecular Number (Tb.N), it counts the number of trabecular per unit length and is calculated by:

$$\text{Tb.N} = \frac{1}{\text{Tb.Th} + \text{Tb.Sp}} \quad (1.4)$$

Connectivity density (Conn.D) (Odgaard, 1997) is a topological indicator normalized by total volume describing the bone structure. It is calculated by:

$$\text{Conn.D} = \frac{\beta_1}{TV} \quad (1.5)$$

$$\beta_1 = \beta_0 + \beta_2 - \chi$$

where β_1 is defined as the maximal number of cuts that can be applied to a domain so that it remains connected. For instance, $\beta_1 = 0$ for a line and β_1 for a torus is 2. In this work, β_1 is determined with β_0 , β_2 together with χ . β_0 and β_2 are the number of cavities and connected components. In trabecular micro-architecture analysis, they correspond to bone and marrow particles. χ is the Euler number. The Euler Number is calculated following the work presented in (Toriwaki and Yonekura, 2002).

The trabecular structure is composed of plate-like and rod-like elements. The modification of bone structure is associated to the variation of the number of plate-like and rod-like elements (Hildebrand and RÜEGSEGER, 1997). The structure model index (SMI) was the first parameter proposed to globally assess the plate-like or rod-like nature of object structure. It is defined as:

$$\text{SMI} = \frac{S' \times V}{S^2} \quad (1.6)$$

where S is the surface of 3D object, V is the volume, S' is the deviation of the surface, (change of surface caused by dilation). The SMI of an ideal plate structure is 0, the one of an ideal rod structure is 3, the one of a sphere is 4. For a structure mixed of plates and rods, SMI ranges between 0 to 3. SMI is negative for cavities.

Conversely to the SMI which is a global parameter, a local topological analysis (LTA) method was proposed to identify each voxel as a plate or a rod (Bonnassie et al., 2003). Local topological analysis (LTA) gives access to quantitative evaluation of plate and rod structures. LTA is performed in 3 steps: first, the voxels on the medial axis are labeled (boundary points, nodes, plates and rods), second, all the other voxels in the image are labeled based on the classification of the voxels on the medial axis, third a quantitative analysis is performed on the image after classification (Bonnassie et al., 2003, Peyrin et al., 2010, Pialat et al., 2012). During the third step, the authors have proposed to calculate several bone parameters, such as the ratio of rods, plates and nodes over the total volume (RV/TV, PV/TV, NV/TV) and the ratio of rod over plate (RV/PV). Figure 1.12 shows an illustration of plates/rods classification on the whole trabecular region of a distal radius.

Bone parameters applications

All these proposed bone parameters help researchers to make accurate studies about bone micro architectures and try to relate these parameters to the risk of fracture.

Different bone parameters between genders and ages are studied in (Khosla et al., 2006). They found that young men have a larger BV/TV and Tb.Th than young women (20-29 years), but they have similar Tb.N and Tb.Sp. As time passes by, the Tb.N of women gradually decreases while men have a Tb.N at the same level as before. Even though the thickness of trabecular of men becomes thinner while that of women is thickening, women have a higher risk of fracture because Tb.N is a very determinant factor in resisting and reducing the pressure and the risk of fracture.

Conventionally, BMD is related to the strength of bone. A high BMD may indicates a low risk of fracture. However, this is not the case for Asian women and other racial groups. Asian women has a lower BMD and a lower risk of fracture. This paradox is explained by Walker *et al.* (Walker et al., 2009). They used HR-pQCT to analyze the micro-structure of bone and tried to extract more information which are not reflected by BMD. Finally, they found that Chinese women have more trabecular number and low

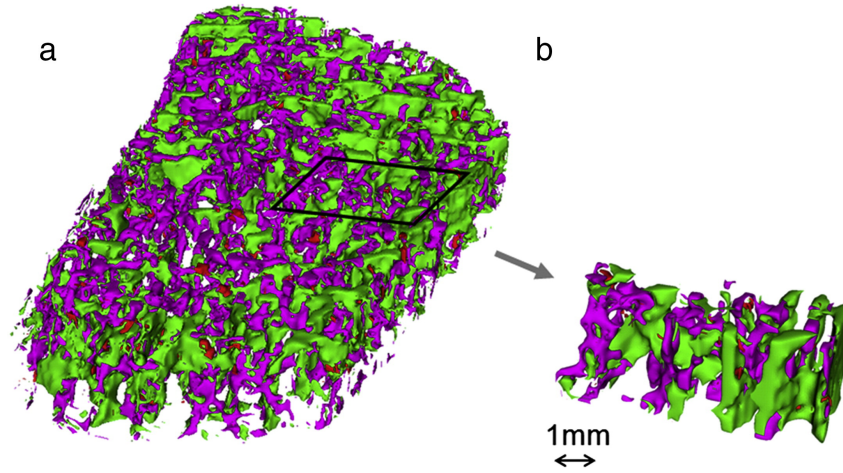


Figure 1.12: a) Results of local topological analysis (LTA) done on the whole trabecular region of a control patient's distal radius. The following color code is used for illustration: pink: pixels of rod-like trabeculae; green: pixels of plate-like trabeculae. b) A portion of the trabecular bone was extracted to better display the plate-like and rod-like trabeculae. Image from (Peyrin *et al.*, 2010)

trabecular separation and inhomogeneity. Such character may enhance the strength of bone and thus the paradox is explained.

Similarly, Lawrence *et al.* in (Riggs and Melton, 1983) summarized two syndromes of osteoporosis: a small subset of postmenopausal women aged from 51 to 65 years suffers from osteoporosis, which is less frequently observed for men at the comparable age; osteoporosis occurs to a large proportion of women and men after 75 years of age. They concluded that the former osteoporosis is due to accelerated disproportionate loss of trabecular bone, while the latter one is caused by proportional loss of cortical and trabecular structure. Geusens *et al.* and Nishiyama *et al.* in (Geusens *et al.*, 2014, Nishiyama and Shane, 2013) present the applications of HR-pQCT in clinical research: observing the structure of normal and diseased bone, comparing the bone micro-architecture among different population, detecting if patients suffered from fractures, monitoring the process of fracture healing, studying the joint diseases. Weaver *et al.* in (Weaver *et al.*, 2016) reveal the effects of diets and physic exercises to the variance of bone mass and micro-architecture. Pialat *et al.* in (Pialat *et al.*, 2012) reveals that the conversion from plate-like structure to rod-like structure implies that the bone turns to be fragile. And the low value of NV/TV indicates that the trabecular network is lacking of connections, i.e. the bone structure is fragile.

1.5 Pros and cons of HR-pQCT images

Even though μ -CT or SR- μ -CT provides CT images with a very good resolution, they are constrained to *ex vivo* imaging. HR-pQCT gives an access to *in vivo* imaging at distal peripheral sites: tibias and radius. Radius is one of the most common fracture sites of human. Tibias, though is not on the list of the common fracture sites, it serves as weight-bearing. HR-pQCT allows for *in vivo* bone strength quantification as well as bone

fracture prediction. Meanwhile, the application of HR-pQCT in clinical research helps to compare the effects of distinct osteoporosis therapies (Cheung et al., 2013).

The accuracy of HR-pQCT quantification is very sensitive to the images resolution. Cheung *et al.* in (Cheung et al., 2013) declared that the density-based measurement, such as Conn.D, may be influenced by scatter and beam hardening artifacts, the criterion SMI is very susceptible to image resolution. Tjong *et al.* in (Tjong et al., 2012) investigates the effects of different voxel size of HR-pQCT by comparing them with μ -CT reference images. They conclude that the $41\mu\text{m}$ voxel size has the best correlation while the $123\mu\text{m}$ voxel size proves the most inaccurate measurements for bone qualification. As far as we know, for human *in vivo* imaging, the voxel size of the first generation HR-pQCT images remains $82\mu\text{m}$. Its spatial resolution at central field is around $130\mu\text{m}$, and decreases to $140\text{--}160\mu\text{m}$ at other regions (Tjong et al., 2012). In fact, as illustrated in Fig.1.3(D), nearly 45% of voxels in a normal trabecular bone are in trabeculae with thickness below $130\mu\text{m}$. As a consequence, measurement errors of HR-pQCT images are visibly introduced due to the resolution of images.

In this thesis, we attempt to improve the resolution of HR-pQCT *in vivo* image so that the measurements of the resolved solution approximate that of the gold standard μ -CT images. In previous works of the group, this problem was addressed by developing total variation based super resolution methods which were successfully evaluated on simulated data of bone micro architecture (PhD A. Toma). However, the application of the method to experimental HR-pQCT was not convincing. In the present work, we will then focus on the development of super resolution methods adapted to such experimental data. Our dataset is composed of high and low resolution image pairs which are scanned from μ -CT and the first generation HR-pQCT on the same samples. The initial problem is to reconstruct super resolution images at $41\mu\text{m}$ voxel size, i.e. improve the resolution by factor 2. However, in the last chapter, we'll see that it is possible to increase this factor to get a voxel size of $24\mu\text{m}$ and thus the image quality is further improved.

We acknowledge that the recent second generation HR-pQCT is capable to scan human *in vivo* at $61\mu\text{m}$ voxel size and gives good measurements (Manske et al., 2015). The resolution of the second generation HR-pQCT images is $95.2\mu\text{m}$ (Manske et al., 2015). Nearly 16% of voxels in a normal trabecular bone are in trabeculae with thickness less than $80\mu\text{m}$, as inferred from Fig.1.3(D). This percentage can be even higher for osteoporotic or suspected osteoporotic bone micro-architecture. To sum up, the resolution of the images provided by the second generation HR-pQCT is still insufficient.

In this thesis, we work with the low resolution HR-pQCT images scanned from the first generalization HR-pQCT. Here are the two reasons:

1. Improve the accuracy of the first generation HR-pQCT measurements which are already employed in practice and in clinical research.
2. We try to determine a super resolution image of $41\mu\text{m}$ (or $24\mu\text{m}$) voxel size which has a better resolution than the images obtained from the second generation HR-pQCT images.

1.6 Summary

To summarize, in this chapter, we introduced the background of osteoporosis and clarify why it's necessary to study trabecular bone micro-architecture. Then we briefly presented some CT devices and bone parameters which are used in the investigation of trabeculae micro structure. Afterwards, we summarized the pros and cons of HR-pQCT and demonstrated our motivations to work on super resolution approaches for this type of images.

Before presenting the approaches applied for our super resolution problem, we present the state of the art of super resolution problem and some notions about mathematic optimization in the second chapter.

Chapter 2

Introduction to super resolution methods and mathematical foundations

As mentioned in the first chapter, the aim of this thesis is to solve a joint super resolution/segmentation problem. The objective of this chapter is to give an overview of super resolution techniques and to introduce some mathematical concepts which are used later. We're going to detail the following aspects successively:

- The inverse problem model for super resolution
- Different approaches for super resolution
- Convex optimization
- Nonconvex optimization
- General introduction of deep learning
- Discussion

2.1 The inverse problem model for super resolution

Let $\mathcal{A}$ be an operator which may be linear or nonlinear between two spaces $\mathcal{X}$ and $\mathcal{Y}$, $\mathcal{A} : \mathcal{X} \rightarrow \mathcal{Y}$. For a given $\mathbf{f}^* \in \mathcal{X}$, we observe an output $\mathbf{g} \in \mathcal{Y}$, and the linear forward problem is formulated as:

$$\mathbf{g} = \mathcal{A}\mathbf{f}^* \quad (2.1)$$

The spaces may be finite dimensional spaces, infinite dimensional Hilbert spaces as L_2 or Banach spaces. In real observation systems, noisy data $\mathbf{g}$ is what we observe, and we usually attempt to determine $\mathbf{f} \in \mathcal{X}$ so that $\mathbf{f}$ approximates the ground truth $\mathbf{f}^*$ as much as possible. Such problems are named as inverse problems. Inverse problems can be defined over topological spaces, finite dimensional spaces, Hilbert spaces or Banach spaces. A well-posed inverse problem satisfies three conditions:

1. The solution exists
2. The solution is unique

3. The operator $\mathcal{A}^{-1}$ is continuous (stability)

Mathematically, the third condition is described as: for all neighborhood $\mathcal{O}(\mathbf{f}^*) \in \mathcal{X}$ of the solution $\mathbf{f}^*$ to Eq.2.1, there exists a neighborhood $\mathcal{O}(\mathbf{g}) \in \mathcal{Y}$ such that for all $\mathbf{g}_\delta \in \mathcal{O}(\mathbf{g})$, $\mathcal{A}^{-1}\mathbf{g}_\delta = \mathbf{f}_\delta$ belongs to $\mathcal{O}(\mathbf{f}^*) \in \mathcal{X}$.

If at least one of the three conditions is not satisfied, then the problem turns to be an ill-posed inverse problem (Kabanikhin, 2008). For some types of operators, the approximated solution depends heavily on noise. This sort of problems are ill-posed, and their solutions have to be stabilized with respect to the noise. In image processing, there are various tasks that can be formulated as inverse problems, and most of them are ill-posed, such as super resolution, deconvolution, image denoising or image reconstruction in CT or MRI for example. In most of them, the operator $\mathcal{A}$ is linear.

We will be particularly interested in image super resolution. More precisely, if we denote $\mathbf{f}^* \in \mathbb{R}^{N'}$ the ground truth, $\mathbf{g} \in \mathbb{R}^N$ the observed data, $\mathcal{A}$ the degrading operator including blurring and down sampling operator, $\mathbf{n}$ some additive noise, then the forward model can be described as:

$$\mathbf{g} = \mathcal{A}\mathbf{f}^* + \mathbf{n} \quad (2.2)$$

The integer N and N' are the discretization levels of the low and high resolution images respectively. If the upsampling factor r is isotropic, for 2D image super resolution, $N' = r^2N$; for 3D image super resolution, $N' = r^3N$. Our objective is to retrieve a resolved image $\mathbf{f} \in \mathbb{R}^{N'}$ from low resolution image $\mathbf{g}$ so that $\mathbf{f}$ is close to $\mathbf{f}^*$.

2.2 Different approaches for super resolution

Super resolution is an inverse problem which has tremendous applications in astronomy, medical imaging, artificial intelligence, industry, etc. Multi-frames super resolution determines super resolution images from a number of low resolution images. Single image resolution improves image resolution based on one single low resolution image. Generally speaking, the super resolution problem may be solved by four approaches:

- Interpolation in spatial domain
- Super resolution in the frequency or wavelet domain domain
- Regularization approaches
- Machine learning

2.2.1 Interpolation in spatial domain

Spatial interpolation estimates the interpolated point by taking account of its neighbors. One of the most simple case is the nearest neighbor interpolation. The interpolated point is equal to the value of its nearest neighbors. Here we are going to give a brief introduction about spatial interpolation methods, such as polynomial interpolation, Bézier curves and Spline interpolation.

General polynomial interpolation

Linear interpolation for two points of value c_0 and c_1 at x_0 and x_1 , with $x_0 < x_1 \in \mathbb{R}$, the interpolating function $\mathcal{F}$ is written as

$$\mathcal{F}(x|c_0, c_1; x_0, x_1) = \frac{x_1 - x}{x_1 - x_0}c_0 + \frac{x - x_0}{x_1 - x_0}c_1 \quad (2.3)$$

as a result, $\mathcal{F}(x)$ is a straight line fixed by two points. If $x \in [x_0, x_1]$, $\mathcal{F}$ is an interpolation operator, otherwise it's an extrapolation operator. For 3 points interpolation, $x_0 < x_1 < x_2$, the function is written as

$$\mathcal{F}(x|c_0, c_1, c_2; x_0, x_1, x_2) = \frac{x_2 - x}{x_2 - x_0}\mathcal{F}(x|c_0, c_1; x_0, x_1) + \frac{x - x_0}{x_2 - x_0}\mathcal{F}(x|c_1, c_2; x_1, x_2) \quad (2.4)$$

this equation can be extended for more points interpolation. Notice that in Eq.2.4, $\mathcal{F}(x|c_0, c_1, c_2; x_0, x_1, x_2)$ is a convex combination of $\mathcal{F}(x|c_0, c_1; x_0, x_1)$ and $\mathcal{F}(x|c_1, c_2; x_1, x_2)$ with $x \in [x_0, x_2]$. While $\mathcal{F}(x|c_0, c_1; x_0, x_1)$ is supposed to be a convex combination of c_0, c_1 with $x \in [x_0, x_1]$, and similarly for $\mathcal{F}(x|c_1, c_2; x_1, x_2)$.

Bézier curves

Bézier curve has been proposed by a French engineer Pierre Bézier (1910-1999). The final estimation function $\mathcal{F}(x)$ is a convex combination of $\{c_i = \mathcal{F}(x_i)\}_{i=0}^{i=2}$ (more points are also allowed). For 3 interpolation points, the interpolation function is defined as

$$\begin{aligned} \mathcal{F}_b(x|c_0, c_1; x_0, x_1) &= (1 - x)c_0 + xc_1 \\ \mathcal{F}_b(x|c_0, c_1, c_2; x_0, x_1, x_2) &= (1 - x)\mathcal{F}_b(x|c_0, c_1; x_0, x_1) + x\mathcal{F}_b(x|c_1, c_2; x_1, x_2) \\ &= (1 - x)^2c_0 + 2x(1 - x)c_1 + x^2c_2 \end{aligned} \quad (2.5)$$

where $x \in [0, 1]$. Eq.2.5 has a good property that the sum of the coefficients of $\{c_i\}_{i=0}^{i=2}$ is 1. In fact, the linear combination with weights $1 - x$ and x follows the distribution of binomial expansion, if there are n points to interpolate, $\{c_i\}_{i=0}^{i=n}$, then their correspondent weights are $C_n^i x^i (1 - x)^{n-i}$ and fulfill $\sum_i C_n^i x^i (1 - x)^{n-i} = 1$. Therefore, we obtain a convex combination for n points, which is not the case for general polynomial interpolation. In Bézier curve, $\{c_i\}_{i=0}^{i=n}$ are called control points, $\mathcal{F}_b(x)$ in Eq.2.5 is called Bernstein polynomials, $\{x_i\}_{i=0}^{i=n}$ are knots or knot vector.

When n is very large, the complexity of Bézier curve increases. The solution of this problem is to decompose n control points into several sets, then implement the interpolation within subset so that the complexity is decreased and finally connect these segments. To keep the continuity of the curve at the joint points, it is necessary to set the difference of the last two control points at previous segment equal to the difference of first two control points at the following segment.

Spline interpolation

The continuity constraint for piecewise connection in Bézier curve is quite strict, it needs to search control points satisfying the requirement. Compared with Bézier curve, spline can joint the segment pieces without considering such constraints.

The interpolation error for polynomials interpolation grows with the number of interpolation points and the degree of the polynomial. Large oscillations are obtained even for a very simple function. The idea of spline interpolation is to interpolate with polynomials functions of degree k and of class C^{k-1} on each sub interval $[x_i, x_{i+1}]$. Such functions are called spline functions. The B-splines function can be obtained with recursive relations. In the variational approach, a spline is obtained as a best interpolant, or the function with the smallest $k - 1^{th}$ derivative among all those matching prescribed function values at certain sites.

Take three interpolation points as an example, instead of 3 knots $\{x_i\}_{i=0}^{i=2}$ in general polynomial interpolation, spline uses 4. To be precise, denote $\mathcal{F}_s(x)$ as the interpolated curve with spline:

$$\begin{aligned}\mathcal{F}_s(x|c_0, c_1; x_0, x_1) &= \frac{x_1 - x}{x_1 - x_0}c_0 + \frac{x - x_0}{x_1 - x_0}c_1, & x \in [x_0, x_1] \\ \mathcal{F}_s(x|c_0, c_1, c_2; x_0, x_1, x_2, x_3) &= \frac{x_2 - x}{x_2 - x_1}\mathcal{F}_s(x|c_0, c_1; x_0, x_2) + \frac{x - x_1}{x_2 - x_1}\mathcal{F}_s(x|c_1, c_2; x_1, x_3), \\ & & x \in [x_1, x_2]\end{aligned}\tag{2.6}$$

By introducing an additional knot, $x \in [x_1, x_2]$ gets rid of the nonconvex combination in general polynomial interpolation. Define the piecewise constant function $\mathcal{B}_i(x)$,

$$\mathcal{B}_i(x) = \begin{cases} 1, & x_i \leq x \leq x_{i+1} \\ 0, & \text{otherwise} \end{cases}$$

Then for 3 points interpolation, B-Spline is formulated as

$$\begin{aligned}\mathcal{F}_{\text{B-Spline}}(x) &= B_0(x) \frac{x_2 - x}{x_2 - x_1} \mathcal{F}_s(x|c_0, c_1; x_0, x_2) + \\ & B_1(x) \frac{x - x_1}{x_2 - x_1} \mathcal{F}_s(x|c_1, c_2; x_1, x_3)\end{aligned}\tag{2.7}$$

Such linear combination leads to a smooth curve and could be extended to a high order. The term $\mathcal{B}_i(x)$ in $\mathcal{F}_{\text{B-Spline}}(x)$ can serve as basis, and in the later segment connection, the reconstruction could be simply generated by a linear combination of these basis. B-spline has many interesting properties, however, they are out of scope of this thesis.

A wide range of techniques have been developed to improve linear image interpolators. Directional image interpolations take advantage of the geometric regularity of image structures by performing the interpolation in a chosen direction along which the image is locally regular. The main difficulty is to locally identify this direction of regularity. Along an edge, the interpolation direction should be parallel to the edge. Many adaptive interpolations have been thus developed with edge detectors ([Shi and Ward, 2002](#)).

2.2.2 Super resolution in frequency-wavelet domain

Apart from spatial interpolation, transforming an image to a different domain is another approach to solve the super resolution problem. Transformation could be based on discrete cosine transform, Fourier transform, discrete wavelet transform. The basic idea is to improve the resolution of image by manipulating the transform coefficients.

Fourier based methods make use of the aliasing property of LR images to reconstruct a high resolution image. Due to their global nature, these methods allow only linear space invariant point spread function (PSF). Moreover, it is difficult to identify a global frequency domain as an priori knowledge to overcome ill-posedness.

Tsai *et al.* are the first researchers who dealt with multi-frame super resolution problems in frequency domain (Tsai, 1984). With a sequence of low resolution images of the same scene, they assumed that the only variance between these images was shifting. The work of Tsai *et al.* exploit the relationship between the continuous Fourier transform of the unknown high resolution image and the discrete Fourier Transform of the low resolution image. In this paper, they assumed that the sampling operator is an impulse response, therefore no over alias was considered, neither was noise or blurring effects. Nevertheless, this work has been a landmark for super resolution research in the frequency domain.

Another method mentioned in (Borman and Stevenson, 1998) took alias and motion into account for super resolution. The main idea is to first estimate the motion, then reverse the transformation on the discrete Fourier transform of the observed data. A number of methods in the literature could estimate the motion or rotation from multiple images. (Kim and Su, 1993, Reddy and Chatterji, 1996, Marcel et al., 1997, Stone et al., 2001, Lucchese and Cortelazzo, 2000, Gluckman, 2003). Robinson *et al* applied combined Fourier-wavelet deconvolution and denoising for super resolution in (Robinson et al., 2010).

Recently many algorithms have been proposed for super resolution in the wavelet domain (Wang, 1990, Nguyen and Milanfar, 2000, Anbarjafari and Demirel, 2010, Ji and Fermüller, 2006). Carey *at al.* have estimated the unknown details of wavelet coefficients in order to improve the sharpness of the reconstructed image (Carey et al., 1999). The principle of these methods is to compute fine scale wavelet coefficients by extrapolating larger scale wavelet coefficients (Crouse et al., 1998, Anbarjafari and Demirel, 2010).

2.2.3 Regularization methods for super resolution

Another class of methods to solve super resolution problems is based on regularization. The solution is obtained by minimizing a regularization functional including a data fitting term and a regularization term, corresponding to some a priori information about the solution, such as sparsity or smoothness. In the literature, the regularization functional is written as:

$$\arg \min_f ||\mathcal{A}f - g||_2^2 + \mu \mathcal{R}(f) \quad (2.8)$$

where μ is a regularization parameter, balancing the effects of data fitting term and regularization term.

The data fitting term in Eq.2.8 uses L_2 norm to measure the distance between $\mathcal{A}f$ and g . It corresponds to an assumption of Gaussian noise on the data. L_2 -norm is the most popular distance for data fitting term in the literature, while other distances have been studied as well, such as mutual information (Hermosillo et al., 2002, Li et al., 2017b) or Kullback Leibler divergence (Favati et al., 2010).

The regularization term is crucial for the accuracy of the solution: on the one hand, it integrates prior information into the final result; on the other hand, it reduces over-fitting risks. The Tikhonov regularization is a classical regularization with L_2 norm. It favors

the solution with smallest norm and enhances the smoothness of the solution. Another common used prior is sparsity. For instance, the gradient of piece-wise image is sparse. In this case, the sparsifying transform is the gradient operator. Even though the image has a very complex texture, it could be sparse after being transformed into another domain. Regularization terms have various forms, loosely speaking, they can be divided into two categories: convex and nonconvex. A convex regularization term is easier to minimize than a nonconvex one. However, nonconvex regularizers tend to have a better performance for sparsity enhancement (Candes et al., 2008).

In the following, we present classic regularization norms, more advanced regularization norms and the choice of the regularization parameters.

Classic regularization norms

Tikhonov regularization is based on the square of the L_2 norm $\|\cdot\|_2^2$. This is a convex and smooth regularization term and is easier to derivate than other regularization terms presented in the following. It has been extensively used during several decades, but its drawback is that it can over smooth the solution

Lasso regularization is based on the L_1 norm $\|\cdot\|_1$ (Tibshirani, 1996). Lasso is a convex but non smooth regularization term. Both Tikhonov and Lasso have many advantages, while it's noteworthy that Tikhonov is good at group selection, its solution is less sparse compared to Lasso. Lasso provides sparse solutions, thus not good at group selection. Let us take one example from (Tibshirani, 1996), the data are generated from the model:

$$y = 0.5x_1 + 0.5x_2 \quad (2.9)$$

Hence the mathematical problem is to find β_1 and β_2 such that

$$y = \beta_1x_1 + \beta_2x_2 \quad (2.10)$$

The Tikhonov regularization term is

$$\beta_1^2 + \beta_2^2 \quad (2.11)$$

while the Lasso regularization term is

$$|\beta_1| + |\beta_2| \quad (2.12)$$

Assume having three candidates,

$$\begin{aligned} y &= 0.1x_1 + 0.9x_2 \\ y &= 0.3x_1 + 0.7x_2 \\ y &= 0.5x_1 + 0.5x_2 \end{aligned} \quad (2.13)$$

their Tikhonov regularization can distinguish them by comparing Eq.2.11:

$$\begin{aligned} \text{For the 1}^{st} \text{ candidate: } & 0.1^2 + 0.9^2 = 0.82 \\ \text{For the 2}^{nd} \text{ candidate: } & 0.3^2 + 0.7^2 = 0.58 \\ \text{For the 3}^{rd} \text{ candidate: } & 0.5^2 + 0.5^2 = 0.50 \end{aligned} \quad (2.14)$$

therefore the third candidate is selected. Yet, Lasso can't distinguish the 3 proposed candidates:

$$\begin{aligned} \text{For the 1}^{st} \text{ candidate: } & |0.1| + |0.9| = 1 \\ \text{For the 2}^{nd} \text{ candidate: } & |0.3| + |0.7| = 1 \\ \text{For the 3}^{rd} \text{ candidate: } & |0.5| + |0.5| = 1 \end{aligned} \quad (2.15)$$

This example shows that Tikhonov regularization favors the group selection: both β_1 and β_2 are far from 0 in the solution given by Tikhonov term. However, the three candidates can't be distinguished by Lasso regularization. To show the effects of Lasso, we assume that there are three other candidates for the problem Eq.2.10:

$$\begin{aligned} y &= \sqrt{0.1}x_1 + \sqrt{0.9}x_2 \\ y &= \sqrt{0.3}x_1 + \sqrt{0.7}x_2 \\ y &= \sqrt{0.5}x_1 + \sqrt{0.5}x_2 \end{aligned} \quad (2.16)$$

In this case, the Tikhonov regularization can't distinguish the three candidates, while Lasso tends to choose the first candidate as solution. This result shows that Lasso regularization prefers the solution which is sparse and at least one coefficient near 0.

Figure 2.1 explains the difference between Tikhonov and Lasso in another way. The ellipse contours may represent the data fitting term. The geometrical figures including the original point stand for the different regularization norms. The intersection points (green bold dots) of ellipses and geometrical figures are potential solutions. If the solution is characterized by two components, θ_1 and θ_2 , it is obvious to see that L_1 -norm favors solutions on the axis, i.e. preserves only one component in this example. However, L_2 -norm manages to find a solution with two components. If the L_1 norm is replaced by a norm mixing L_1 and L_2 norm, we obtain a norm called group lasso optimization or elastic net regularization.

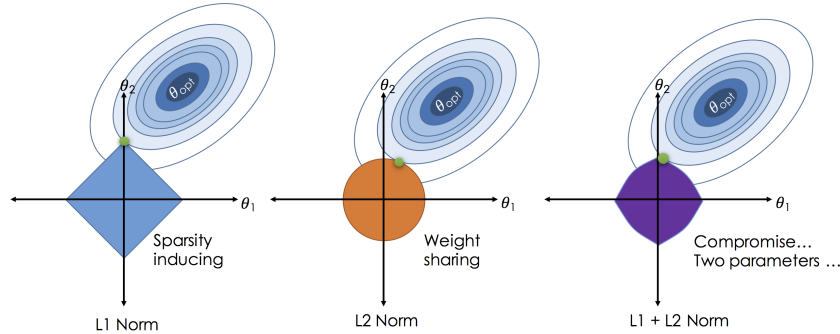


Figure 2.1: Illustration of Tikhonov, Lasso or Elastic net regularization. The image is drawn by Joseph E. Gonzalez. Website: http://www.ds100.org/sp17/assets/notebooks/linear_regression/Regularization.html

Elastic net regularization is a linear combination of Tikhonov and Lasso (Zou and Hastie, 2005). Researchers in (Zou and Hastie, 2005) mentioned that elastic net regularization outperforms Tikhonov and Lasso regularization. Numerically, elastic net regularization term generalizes Tikhonov and Lasso. And Figure 2.1 also shows that elastic net regularization term combines the group selection advantage of Tikhonov and sparse solution of Lasso.

Advanced regularization terms

Total Variation (TV) regularization smooths images by enhancing the sparsity of the image variation. The image variation is defined as the L_1 norm of the gradient (Chambolle et al., 2010):

$$\mathbf{TV}(\mathbf{f}) = \int_{\Omega} |\nabla \mathbf{f}(x)| dx \quad (2.17)$$

This is a non smooth regularization term. It has an adequate performance on image denoising, particularly for the piecewise images. The TV regularization computes the first order derivative of images, higher order derivatives can also be used (Toma et al., 2014b).

The L_0 norm is another widely discussed norm which can be used for sparse signals recovery. The L_0 -norm is defined as:

$$\|\mathbf{f}\|_0 = |\{i, f_i \neq 0\}| \quad (2.18)$$

Loosely speaking, L_0 -norm counts the number of nonzero elements. Minimizing the L_0 -norm is a nonconvex NP-hard problem. Candes et al. work on an approximate of the L_0 -minimization problem which is the log-penalty problem, i.e. the log-penalty sparsity measure is used instead of the L_0 -norm (Candes et al., 2008). For a sparse signal $\mathbf{f}$

$$\mathbf{f} = [f_1, f_2, f_3, \dots, f_n]^T \in \mathbb{R}^{N'} \quad (2.19)$$

the log-sum penalty is written as

$$\min_{\mathbf{f} \in \mathbb{R}^{N'}} \sum_{i=1}^{N'} \log(\epsilon + |f_i|), \quad i = 1, 2, \dots, N' \quad (2.20)$$

where ϵ is a small constant. Candes et al. solve this problem iteratively by

$$\mathbf{f}^{(k+1)} = \arg \min \sum_{i=1}^{N'} \frac{|f_i|}{|f_i^{(k)}| + \epsilon}, \quad i = 1, 2, \dots, N' \quad (2.21)$$

where k is the iteration number. Denote $\boldsymbol{\mu}$ as weights defined as

$$\boldsymbol{\mu} = \text{diag}(\mu_1, \mu_2, \mu_3, \dots, \mu_n) \in \mathbb{R}^{2N'} \quad (2.22)$$

then the algorithm amounts to

$$\mathbf{f}^{(k+1)} = \arg \min_{\mathbf{f}} \|\boldsymbol{\mu}^{(k)} \mathbf{f}\|_1, \quad \text{with} \quad \mu_i^{(k)} = \frac{1}{|f_i^{(k)}| + \epsilon} \quad (2.23)$$

The principle is to erase small signals by increasing the weights of the regularization while keeping large signals by reducing their penalty weights.

The Mumford Shah functional is nonconvex functional to recover piecewise constant images (Vese and Chan, 2002). It not only takes account of intensity variation within piece-wise constant, but also the minimize the length of its boundary. Elleuch et al. in (Elleuch et al., 2016) proposed a gradually nonconvex optimization based on L_1 minimization. The main idea is to gradually approximate L_0 norm by iteratively adapting an hyper-parameter.

Choice of regularization parameters

All the previous methods require to choose the regularization parameters which balance the data term and regularization terms. For this purpose, besides traditional sweeping strategy, it's possible to use L-curve (Hansen, 1999) to choose the regularization parameter for Tikhonov regularization or the Morozov principle (Engl et al., 1996). L-curve is a log-log curve of the regularization term $\|\mathbf{f}\|_2$ versus data fitting term norm $\|\mathbf{A}\mathbf{f} - \mathbf{g}\|_2$. The points on the curve correspond to the solutions obtained with different values of the regularization parameter μ . When the curve evolves dramatically, its correspondent μ is assumed to be the most optimal compromise between the data fitting term and the regularization term. Morozov principle chooses parameter μ by referring $\|\mathbf{A}\mathbf{f} - \mathbf{g}\|_2 = \sigma$, where σ denotes the noise level of the observed image (Engl et al., 1996).

The application of regularization in super resolution

Toma *et al.* have investigated the application of TV and higher order TV in (Toma et al., 2014b, Toma et al., 2014a). Tofighi *et al.* considered to add constraints in spatial domain and Fourier domain (Tofighi et al., 2016). In the application of dictionary learning in super resolution, regularization term is also used to enhance the sparsity of the representation coefficients (Jing et al., 2015). In the application of deep learning in super resolution, the regularization term is used to avoid overfitting (LeCun et al., 2015).

2.2.4 Dictionary learning

As previously mentioned, regularization is a way to inject priors into solutions. Sometimes, regularizers penalize images directly, like TV, sometimes indirectly, e.g. dictionary learning. Dictionary learning reconstructs images via sparse representation over a basis of dictionary atoms.

This section is devoted to the presentation of the application of dictionary learning in super resolution problems. It first explains how dictionary learning works, then reviews different approaches of super resolution problems and the interest of sparsity in dictionary learning approaches.

Dictionary learning introduction

In dictionary learning, the set of basis vectors is named as dictionary, and each basis vector is named as atom. Atoms can be defined by a set of known basis functions, such as steerable filters, or be learned from training samples. Let N_f be the dimension of one image patch, the equation to represent a signal $\mathbf{f}_0 \in \mathbb{R}^{N_f \times 1}$ over the dictionary $\Phi \in \mathbb{R}^{N_f \times p}$ with coefficients $\alpha_0 \in \mathbb{R}^{p \times 1}$ is written as:

$$\mathbf{f}_0 = \Phi \alpha_0 \quad (2.24)$$

where $\Phi = \{\mathbf{d}_1, \mathbf{d}_2, \dots, \mathbf{d}_p\}$, $\mathbf{d}_i \in \mathbb{R}^{N_f \times 1}$ is one atom in the dictionary, $i = 1, 2, \dots, p$. The first step in dictionary learning is to learn the dictionary, then reconstruct images with the learned dictionary.

K-SVD (Aharon et al., 2006) is one of the most classic dictionary learning approaches. We consider a training dataset $\mathbf{f}_{\text{train}}$ with m signals, its sparse coefficient over the dictionary Φ denoted as α , with

$$\begin{aligned}\mathbf{f}_{\text{train}} &= \{\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_m\} && \in \mathbb{R}^{N_f \times m} \\ \alpha &= \{\alpha_1, \alpha_2, \dots, \alpha_m\} && \in \mathbb{R}^{p \times m} \\ \Phi &= \{\mathbf{d}_1, \mathbf{d}_2, \dots, \mathbf{d}_p\} && \in \mathbb{R}^{N_f \times p}\end{aligned} \quad (2.25)$$

the optimization functional is defined as

$$\arg \min_{\Phi, \alpha} \|\mathbf{f}_{\text{train}} - \Phi \alpha\|_F^2 \quad \text{s.t.} \quad \|\alpha_i\|_0 \leq p', \quad \forall i = 1, 2, \dots, m. \quad (2.26)$$

This equation picks up p' atoms for the representation of each signal $\mathbf{f}_i$, with $p' \ll p$. More precisely, it initializes the dictionary by randomly picking up p atoms from the training dataset. Then it determines the coefficients of each sample in the training dataset. With these temporarily estimated coefficients, atoms will be updated with singular value decomposition (SVD) sequentially. When an atom is updated, its relevant coefficients are updated as well. Updating the atoms via SVD not only enhances the sparsity of coefficients, but also normalizes the atoms which makes the algorithm more stable.

In reconstruction stage, the degraded image can be sparsely represented by optimizing the functional in Eq. 2.26 so that structural information could be well preserved while noise is filtered away.

Application of dictionary learning in super resolution problem

A joint dictionary learning method for super resolution has been studied by Yang *et al.* (Yang et al., 2010a) in 2010. They proposed to learn dictionary in a joint way. During the learning stage, the optimization function is written as:

$$\begin{aligned}\min_{\Phi^f, \Phi^g, \alpha} & \frac{1}{N_f} \|\mathbf{f}_{\text{train}} - \Phi^f \alpha\|_F^2 + \frac{1}{N_g} \|\mathbf{g}_{\text{train}} - \Phi^g \alpha\|_F^2 + \mu \left(\frac{1}{N_f} + \frac{1}{N_g} \right) \|\alpha\|_1 \\ \text{s.t.} & \quad \|\mathbf{d}_i^f\|_2 \leq 1, \|\mathbf{d}_i^g\|_2 \leq 1, \quad i = 1, 2, \dots, p\end{aligned} \quad (2.27)$$

where $\mathbf{f}_{\text{train}}$ and $\mathbf{g}_{\text{train}}$ are high and low resolution images training set, N_f and N_g are the dimensions of one patch in vector form in $\mathbf{f}_{\text{train}}$ and $\mathbf{g}_{\text{train}}$, Φ^f and Φ^g are dictionaries of high and low resolution images, α is the sparse coefficient of $\mathbf{f}_{\text{train}}$ over dictionary Φ^f and $\mathbf{g}_{\text{train}}$ over dictionary Φ^g .

$$\begin{aligned}\mathbf{f}_{\text{train}} &= \{\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_m\} && \in \mathbb{R}^{N_f \times m} \\ \mathbf{g}_{\text{train}} &= \{\mathbf{g}_1, \mathbf{g}_2, \dots, \mathbf{g}_m\} && \in \mathbb{R}^{N_g \times m} \\ \alpha &= \{\alpha_1, \alpha_2, \dots, \alpha_m\} && \in \mathbb{R}^{p \times m} \\ \Phi^f &= \{\mathbf{d}_1^f, \mathbf{d}_2^f, \dots, \mathbf{d}_p^f\} && \in \mathbb{R}^{N_f \times p} \\ \Phi^g &= \{\mathbf{d}_1^g, \mathbf{d}_2^g, \dots, \mathbf{d}_p^g\} && \in \mathbb{R}^{N_g \times p}\end{aligned} \quad (2.28)$$

In (Yang et al., 2010a), Φ^f and Φ^g are concatenated, and also $\mathbf{f}_{\text{train}}$ and $\mathbf{g}_{\text{train}}$. The optimization function turns to be:

$$\min_{\Phi^f, \Phi^g, \alpha} = \left\| \left[\frac{1}{\sqrt{N_f}} \mathbf{f}_{\text{train}}^T, \frac{1}{\sqrt{N_g}} \mathbf{g}_{\text{train}}^T \right]^T - \left[\frac{1}{\sqrt{N_f}} (\Phi^f)^T, \frac{1}{\sqrt{N_g}} (\Phi^g)^T \right]^T \alpha \right\|_F^2 + \mu' \|\alpha\|_1 \quad (2.29)$$

During the super resolution stage, the function involved is:

$$\alpha' = \arg \min_{\alpha} \|\mathbf{g}_{\text{test}} - \Phi^g \alpha\|_F^2 + \mu \|\alpha\|_1 \quad (2.30)$$

where $\mathbf{g}_{\text{test}}$ is the test low resolution patches. And the final super resolution patches are estimated by

$$\mathbf{f}_{\text{test}} = \Phi^f \alpha' \quad (2.31)$$

One drawback is that during the training stage, Φ^f and Φ^g are updated jointly. As a result, neither $\mathbf{f}_{\text{train}}$ nor $\mathbf{g}_{\text{train}}$ have been sufficiently expressed over Φ^f and Φ^g .

Based on this observation, the author came up with a coupled-dictionary learning method (Yang et al., 2012) two years later. Coupled-dictionary learning approach couples two data fitting term to update dictionary atoms with the optimization function:

$$\begin{aligned} \min_{\Phi^f, \Phi^g, \alpha} \quad & \rho \|\mathbf{f}_{\text{train}} - \Phi^f \alpha\|_F^2 + (1 - \rho) \|\mathbf{g}_{\text{train}} - \Phi^g \alpha\|_F^2 + \mu \|\alpha\|_1 \\ \text{s.t.} \quad & \|\mathbf{d}_i^f\|_2 \leq 1, \|\mathbf{d}_i^g\|_2 \leq 1 \end{aligned} \quad (2.32)$$

Differing from Equation 2.29, the parameter ρ in Equation 2.32 is a hyper-parameter which is introduced to balance the two data terms. Moreover, the two dictionaries in Equation 2.32 are updated sequentially. This favors the accuracy of the sparse representation, therefore improves the efficiency of the method.

In the coupled dictionary method, the coefficients of sparse representations are shared over both dictionaries. Yet, common coefficients lead to a limited reconstruction quality and may reduce the accuracy over the current dictionary learning. In order to improve the method, Wang *et al.* proposed semi-coupled dictionary learning (Wang et al., 2012)(SCDL). The principle is to find mapping matrix between the sparse coefficients of high and low resolution patches.

Moreover, He *et al.* used a Bayesian method and a Beta process for coupled dictionary learning (He et al., 2013), and compared their method with other super resolution techniques, as shown in Fig.2.2. Authors in (Wang et al., 2012, He et al., 2013) have considered locality constraints, but they did not integrate low rank constraints. Jing *et al.* in (Jing et al., 2015) incorporated low rank in SCDL, and Wang *et al.* have combined encoder with low rank to achieve marginalized denoising dictionary learning with locality constraint (Wang et al., 2018), but not for super resolution problems.

2.3 Convex optimization

In many inverse problems, an approximate solution can be found by convex optimization methods. Some preliminary notions of convex or nonconvex optimization are given in Annexes (Artacho et al., 2014, Gasso et al., 2009). In many inverse problems, an approximate solution can be found by convex optimization methods. Convex optimization is the basis of more complex optimization problems. When a problem can be formulated as a convex model, its optimization solution is global. The minimizer is not necessarily unique unless the problem is strictly convex. Convex optimization problem consists of an



Figure 2.2: Visual comparison of factor of 2 super-resolution results. The upper row shows the SR results of the image Lena. The lower row shows the SR results of the image House. The SCDL and the beta process are two approaches proposed in (Wang et al., 2012) and (He et al., 2013), respectively. The figure is adapted from (He et al., 2013)

objective function $\mathcal{F}(\mathbf{f})$ and a sequence of equality or inequality constraints described by convex functions: $\mathcal{F}_i(\mathbf{f}), i = 1, 2, \dots, m$:

$$\begin{aligned} \min \quad & \mathcal{F}(\mathbf{f}) \\ \text{subject to} \quad & \mathcal{F}_i(\mathbf{f}) \leq 0, \quad i = 1, 2, \dots, m \end{aligned} \quad (2.33)$$

In a convex optimization problem, if the objective function is quadratic and all of its constraints are linear equations, it's possible to find its analytical solution. When the optimization problem includes linear equation constraints and second order differentiable objective functions, we can apply Newton's method to find its solution by an iterative algorithm.

One method to tackle this problem is to gradually approximate the inequality constraint. The iterative solutions, also called as central points, evolve along what is called central path (points set) (Boyd and Vandenberghe, 2004).

Constraints including inequality are generally studied with Lagrangian function. In this thesis, most of our constraints are equalities, nevertheless they may be nonlinear or even nonsmooth. Many approaches have been proposed to solve such problems. And here we present proximal operator and alternate direction method of multiplier which are very useful tools for convex optimization.

2.3.1 Proximal operator

In this subsection, we present the definition of proximal operator and show one of its interpretations.

Proximal operator. $\text{prox}_{\mathcal{F}} : \mathbb{R}^{N'} \rightarrow \mathbb{R}^{N'}$ of $\mathcal{F}$ is defined as

$$\text{prox}_{\mathcal{F}}(\mathbf{v}) = \arg \min_{\mathbf{f}} (\mathcal{F}(\mathbf{f}) + \frac{1}{2} \|\mathbf{f} - \mathbf{v}\|_2^2) \quad (2.34)$$

where $\|\cdot\|_2$ is the Euclidean norm. With scale parameter θ , the proximal operator is defined as

$$\text{prox}_{\theta\mathcal{F}}(\mathbf{v}) = \arg \min_{\mathbf{f}} (\mathcal{F}(\mathbf{f}) + \frac{1}{2\theta} \|\mathbf{f} - \mathbf{v}\|_2^2) \quad (2.35)$$

Scale parameter θ adapts the balance between the effects of two terms. Proximal operator attempts to find a proximal point which minimizes $\mathcal{F}$ while being closest to $\mathbf{v}$. It can be seen that $\mathbf{v}$ is optimizing $\mathcal{F}(\cdot)$ if and only if $\mathbf{v} = \text{prox}_{\mathcal{F}}(\mathbf{v})$.

A generalization of backward Euler method

The relation between proximal operator $\text{prox}_{\mathcal{F}}(\mathbf{f})$ and the subdifferential operator $\partial\mathcal{F}$ is revealed by

$$\text{prox}_{\theta\mathcal{F}}(\cdot) = (\mathbf{I} + \theta\partial\mathcal{F})^{-1}(\cdot) \quad (2.36)$$

It's noteworthy that $\mathcal{F}(\cdot)$ is a convex function, $\mathcal{F}(\cdot) + \frac{1}{2\theta} \|\mathbf{f} - \mathbf{v}\|_2^2$ is strongly convex, thus strictly convex, therefore the minimizer is global and unique. As a result, $(\mathbf{I} + \theta\partial\mathcal{F})^{-1}$ is single-valued, even though $\partial\mathcal{F}$ is not. According to the definition in Eq.2.35, $\hat{\mathbf{f}} = \text{prox}_{\theta\mathcal{F}}(\mathbf{v})$, then

$$\begin{aligned} \hat{\mathbf{f}} &= (\mathbf{I} + \theta\partial\mathcal{F})^{-1}(\mathbf{v}) \\ \mathbf{v} &= \hat{\mathbf{f}} + \theta\partial\mathcal{F}(\hat{\mathbf{f}}) \\ \mathbf{0} &= \frac{1}{\theta}(\hat{\mathbf{f}} - \mathbf{v}) + \partial\mathcal{F}(\hat{\mathbf{f}}) \end{aligned} \quad (2.37)$$

The third step in Eq.2.37 corresponds to the subdifferential of $\mathcal{F}(\mathbf{f}) + \frac{1}{2\theta} \|\mathbf{f} - \mathbf{v}\|_2^2$ with respect to $\mathbf{f}$ when $\mathbf{f} = \hat{\mathbf{f}}$ in Eq.2.35. Further more, if $\mathcal{F}$ is differentiable, it reduces to the backward Euler method with discretization, which is written as:

$$\frac{\mathbf{f}^{k+1} - \mathbf{f}^k}{\theta} = -\nabla\mathcal{F}(\mathbf{f}^{k+1}) \quad (2.38)$$

where $\mathbf{f}^k$ is the intermediate solution at k^{th} iteration. Thus the proximal operator can be regarded as backward Euler method but generalizes gradient descent methods for non differentiable functions (Parikh et al., 2014, Combettes and Pesquet, 2011). Even if $\mathcal{F}$ is not smooth, the proximal operator leads to a smooth function.

Forward Euler method with discretization is defined as

$$\frac{\mathbf{f}^{k+1} - \mathbf{f}^k}{\theta} = -\nabla\mathcal{F}(\mathbf{f}^k) \quad (2.39)$$

where $\nabla\mathcal{F}(\mathbf{f}^k)$ is explicitly calculated. On the contrary, backward Euler method updates variables implicitly.

2.3.2 Alternate direction method of multiplier(ADMM)

ADMM is one of the most popular algorithms to solve a convex optimization problem. It can be formulated as a proximal operator algorithm. Here we follow the presentation of ADMM in (Boyd, 2011), which explains the algorithm via primal and dual problem optimization. Consider an optimization problem with a linear constraint:

$$\begin{aligned} \arg \min \mathcal{F}(\mathbf{f}) \\ \text{s.t. } \mathcal{A}\mathbf{f} = \mathbf{b} \end{aligned} \quad (2.40)$$

where $\mathcal{F}(\mathbf{f})$ is a convex function. Then Lagrangian functional is formulated as

$$\mathcal{L}(\mathbf{f}, \boldsymbol{\lambda}) = \mathcal{F}(\mathbf{f}) + \boldsymbol{\lambda}^T(\mathcal{A}\mathbf{f} - \mathbf{b}) \quad (2.41)$$

so that we can define primal and dual problems:

$$\begin{aligned} \text{primal problem: } \min_{\mathbf{f}} \mathcal{L}(\mathbf{f}, \boldsymbol{\lambda}) \\ \text{dual problem: } \max_{\boldsymbol{\lambda}} \mathcal{L}(\mathbf{f}, \boldsymbol{\lambda}) \end{aligned} \quad (2.42)$$

$\boldsymbol{\lambda}$ is updated with gradient method, hence

$$\begin{aligned} \mathbf{f}^{k+1} &= \arg \min_{\mathbf{f}} \mathcal{L}(\mathbf{f}, \boldsymbol{\lambda}^k) \\ \boldsymbol{\lambda}^{k+1} &= \boldsymbol{\lambda}^k + \mu(\mathcal{A}\mathbf{f}^{k+1} - \mathbf{b}) \end{aligned} \quad (2.43)$$

The algorithm mentioned above is named as **dual ascent**. When an optimization problem has n variables and can be modeled as a separable function

$$\mathcal{F}(\mathbf{f}_1, \mathbf{f}_2, \dots, \mathbf{f}_n) = \sum_i^n \mathcal{F}_i(\mathbf{f}_i) \quad (2.44)$$

where n is the number of variables. We get:

$$\begin{aligned} \mathbf{f}_i^{k+1} &= \arg \min_{\mathbf{f}_i} \mathcal{L}_i(\mathbf{f}_i, \boldsymbol{\lambda}^k) \\ \boldsymbol{\lambda}^{k+1} &= \boldsymbol{\lambda}^k + \mu(\mathcal{A} \sum_i^n \mathbf{f}_i^{k+1} - \mathbf{b}) \end{aligned} \quad (2.45)$$

It's noteworthy that variables $\mathbf{f}_i$ are updated in parallel, then $\boldsymbol{\lambda}_k$ is updated.

However, this Lagrangian functional can be applied only for differentiable functions. In order to make the algorithm be feasible for more optimization scenarios, a quadratic term has been introduced so that algorithm still works even though $\mathcal{F}$ is non-differentiable. This new method is named as Augmented Lagrangian method. Consider the problem in 2.40, its augmented Lagrangian functional $\mathcal{L}_A$ is written as

$$\mathcal{L}_A(\mathbf{f}, \boldsymbol{\lambda}) = \mathcal{F}(\mathbf{f}) + \boldsymbol{\lambda}^T(\mathcal{A}\mathbf{f} - \mathbf{b}) + \frac{\rho}{2} \|\mathcal{A}\mathbf{f} - \mathbf{b}\|_2^2 \quad (2.46)$$

where ρ is a given positive number. Then the variables are updated by

$$\begin{aligned} \mathbf{f}_i^{k+1} &= \arg \min_{\mathbf{f}_i} \mathcal{L}_i(\mathbf{f}_i, \boldsymbol{\lambda}^k) \\ \boldsymbol{\lambda}^{k+1} &= \boldsymbol{\lambda}^k + \rho(\mathcal{A} \sum_i^n \mathbf{f}_i^{k+1} - \mathbf{b}) \end{aligned} \quad (2.47)$$

This algorithm is named as Method of Multiplier (MM). Despite its robustness, $\mathcal{L}_A$ can't update variables in parallel because of the quadratic term. For instance, consider to solve problem

$$\begin{aligned} \arg \min_{\mathbf{f}} \mathcal{F}(\mathbf{f}) &= \arg \min_{\mathbf{f}_1, \mathbf{f}_2} \mathcal{F}_1(\mathbf{f}_1) + \mathcal{F}_2(\mathbf{f}_2) \\ \text{s.t. } \mathcal{A}\mathbf{f}_1 + \mathcal{B}\mathbf{f}_2 &= \mathbf{b} \end{aligned} \quad (2.48)$$

Its augmented Lagrangian functional is formulated as

$$\mathcal{L}_A(\mathbf{f}_1, \mathbf{f}_2, \boldsymbol{\lambda}) = \mathcal{F}_1(\mathbf{f}_1) + \mathcal{F}_2(\mathbf{f}_2) + \boldsymbol{\lambda}^T(\mathcal{A}\mathbf{f}_1 + \mathcal{B}\mathbf{f}_2 - \mathbf{b}) + \frac{\rho}{2} \|\mathcal{A}\mathbf{f}_1 + \mathcal{B}\mathbf{f}_2 - \mathbf{b}\|_2^2 \quad (2.49)$$

Its variables $\mathbf{f}_1$ and $\mathbf{f}_2$ can't be completely separated thus it's not accurate to update them in parallel. Gabay *et al.* in 1976 proposed a framework to solve this problem: alternating direction method of multipliers (ADMM) (Gabay and Mercier, 1976). Its principle is to sequentially update variables. More precisely, it updates only one variable while keeping others fixed.

$$\begin{aligned} \mathbf{f}_1^{k+1} &= \arg \min_{\mathbf{f}_1} \mathcal{L}_A(\mathbf{f}_1, \mathbf{f}_2^k, \boldsymbol{\lambda}^k) \\ \mathbf{f}_2^{k+1} &= \arg \min_{\mathbf{f}_2} \mathcal{L}_A(\mathbf{f}_1^{k+1}, \mathbf{f}_2, \boldsymbol{\lambda}^k) \\ \boldsymbol{\lambda}^{k+1} &= \boldsymbol{\lambda}^k + \rho(\mathcal{A}\mathbf{f}_1^{k+1} + \mathcal{B}\mathbf{f}_2^{k+1} - \mathbf{b}) \end{aligned} \quad (2.50)$$

To sum up the evolution of ADMM: it starts from Lagrangian optimization functional which is constraint to differentiable functions. Then the augmented Lagrangian functional breaks this bottleneck by introducing a quadratic term and its corresponding algorithm is named as method of multiplier. However, this quadratic term destroys the splitting property when the optimization functional is separable, i.e. variables can't be updated in parallel. As a result, ADMM has been proposed to update variables sequentially.

ADMM has been widely applied in super resolution optimization. It can be used to add non-differentiable constraints to the images (Toma et al., 2014a, Peyrin et al., 2015), or in the deconvolution super resolution problem, it assists to alternatively update the super resolution image and the degradation kernel (Ng et al., 2010, Morin et al., 2012), or in the dictionary learning methods for super resolution, it is applied to derive convolutional filters (Gu et al., 2015), and so forth.

2.4 Nonconvex optimization methods

Convex optimization has an important role in nonconvex optimization problems. It can provide an initialization for global optimization, or serve as a lower bound on the solution of global optimization (Boyd and Vandenberghe, 2004). In this section, we briefly present some nonconvex optimization approaches which solve nonconvex problems: proximal approaches; transforming nonconvex functions into the form of convex function by adding quadratic term so that the function turns to be convex; obtaining convex function by Fenchel conjugate transform. The different methods differ by the additionnal terms they introduce to make the connection with the convex optimization and by the convex tools they use. Moreover, some approaches decompose the cost function as the sum of nonconvex/non differentiable functions or as a difference of convex functions.

2.4.1 Accelerated proximal gradient method

In the framework of the accelerated proximal gradient method, the optimization problem is formulated as (Parikh et al., 2014):

$$\min_{\mathbf{f} \in \mathbb{R}^{N'}} \mathcal{F}(\mathbf{f}) + \mathcal{G}(\mathbf{f}) \quad (2.51)$$

where $\mathcal{F} : \mathbb{R}^{N'} \rightarrow \mathbb{R}$ is a differentiable and possibly nonconvex function, $\mathcal{G} : \mathbb{R}^{N'} \rightarrow \mathbb{R} \cup \{+\infty\}$ is convex but possibly nonsmooth function. Hence the problem in 2.51 combines a nonconvex and nonsmooth optimization. The proximal gradient method can be written as

$$\mathbf{f}^{k+1} := \text{prox}_{\theta^k \mathcal{G}}(\mathbf{f}^k - \theta^k \nabla \mathcal{F}(\mathbf{f}^k)) \quad (2.52)$$

where $\theta^k > 0$ is a step size. $\nabla \mathcal{F}$ is Lipchitz continuous with constant L . The smooth function $\mathcal{F}$ is optimized with gradient descent (forward step), the nonsmooth function $\mathcal{G}$ is minimized by proximal operator (backward step). Thus the algorithm can be regarded as forward-backward splitting approach.

Comparing with proximal gradient method, accelerated proximal gradient method has an additional extrapolation step.

$$\begin{aligned} \mathbf{v}^{k+1} &:= \mathbf{f}^k + \varphi^k(\mathbf{f}^k - \mathbf{f}^{k-1}) \\ \mathbf{f}^{k+1} &:= \text{prox}_{\theta^k \mathcal{G}}(\mathbf{v}^{k+1} - \theta^k \nabla \mathcal{F}(\mathbf{v}^{k+1})) \end{aligned} \quad (2.53)$$

where $\varphi^k \in [0, 1)$ is an extrapolation parameter. Conventionally, set $\varphi^0 = 0$ to eliminate $\mathbf{f}^{-1}$.

Li et al. in (Li and Lin, 2015) refined this scheme so that the convergence speed is further accelerated. In their algorithm, θ^k is adapted according to the quality of the extrapolation points.

2.4.2 Ipiano: inertial proximal approach

Another nonconvex nonsmooth optimization method with proximal operator is presented in (Ochs et al., 2014). Dealing with the same problem as Eq.2.53, their solution is found with

$$\mathbf{f}^{k+1} = (\mathbf{I} + \theta^k \partial \mathcal{G})^{-1}(\mathbf{f}^k - \theta^k \nabla \mathcal{F}(\mathbf{f}^k) + \mu^k(\mathbf{f}^k - \mathbf{f}^{k-1})) \quad (2.54)$$

The first term on the right-hand is a proximal operator, similar to Eq.2.36. Compared with the conventional proximal gradient descent method, Eq.2.54 has one more inertial term, $\mu^k(\mathbf{f}^k - \mathbf{f}^{k-1})$. Polyak pointed out in (Polyak, 1964) that such inertial force accelerates the convergence speed with respect to conventional gradient descent method, and the inertia keeps the algorithm evolving along a good direction towards a stationary point of the functional.

2.4.3 Linearly constrained nonconvex nonsmooth minimization

Artina et al. in (Artina et al., 2013) proposed a linearly constrained nonconvex nonsmooth minimization and applied ADMM to solve nonconvex nonsmooth optimization problem. This approach is a special case of the later proposed Proximal Alternating Linearized

Minimization (PALM) method (Bolte et al., 2014). Given the current iterate $\mathbf{v}$, the principle of the algorithm is to find ω such that $\mathcal{F}_{\omega, \mathbf{v}}(\mathbf{f})$ is ν -strongly convex.

$$\mathcal{F}_{\omega, \mathbf{v}}(\mathbf{f}) := \mathcal{F}(\mathbf{f}) + \omega \|\mathbf{f} - \mathbf{v}\|^2 \quad (2.55)$$

where $\mathbf{v} \in \mathbb{R}^{N'}$. Then the minimization of $\mathcal{F}(\mathbf{f})$ turns to minimize $\mathcal{F}_{\omega, \mathbf{v}}$ iteratively, and $\mathbf{v}$ can be regarded as an intermediate approximation of the current solution $\mathbf{f}$. Precisely, let l be the outer loop iterate, k be the inner loop iterate, L_l be the total number of inner iteration at outer iterate l . With constraint $\mathcal{A}\mathbf{g} = \mathbf{v}$, the algorithm solved with ADMM is described as:

$$\mathbf{f}_{(l,0)} = \mathbf{f}_{l-1} := \mathbf{f}_{(l-1, L_{l-1})} \quad (2.56)$$

$$\boldsymbol{\lambda}_{(l,0)} = \boldsymbol{\lambda}_{l-1} := \boldsymbol{\lambda}_{(l-1, L_{l-1})} \quad (2.57)$$

$$\mathbf{f}_{(l,k)} = \arg \min_{\mathbf{f}} \mathcal{F}_{\omega, \mathbf{f}_{l-1}}(\mathbf{f}) - \boldsymbol{\lambda}_{(l,k-1)}^t (\mathcal{A}\mathbf{f} - \mathbf{v}) + \frac{\beta}{2} \|\mathcal{A}\mathbf{f} - \mathbf{v}\|_2^2 \quad \text{for } k = 1, \dots, L_l \quad (2.58)$$

$$\boldsymbol{\lambda}_{(l,k)} = \boldsymbol{\lambda}_{(l,k-1)} - \beta(\mathcal{A}\mathbf{f} - \mathbf{v}) \quad (2.59)$$

Eq.2.58 and 2.59 are solved iteratively within the inner loop, and the stopping criteria is defined as:

$$(1 + \|\boldsymbol{\lambda}_{l-1}\|) \|\mathcal{A}\mathbf{f}_{(l, L_l)} - \mathbf{v}\| \leq \frac{1}{l^\alpha} \quad (2.60)$$

with $\alpha > 1$

2.4.4 DC programming

Another framework to solve nonconvex problem is DC programming. DC is the abbreviation of "Difference of Convex functions". Gasso *et al.* in (Gasso et al., 2009) summarized the application of nonconvex regularization in sparse signal recovering with DC programming. DC solves nonconvex optimization via Fenchel conjugate transform using the fact that the biconjugate of nonconvex function turns out to be its convex hull. Consider the optimization problem to minimize nonconvex function $\mathcal{F}(\cdot)$

$$\min_{\mathbf{f} \in \mathbb{R}^N} \mathcal{F}(\mathbf{f}) \quad (2.61)$$

With DC programming, it is written as

$$\mathcal{F}(\mathbf{f}) = \mathcal{F}_1(\mathbf{f}) - \mathcal{F}_2(\mathbf{f}) \quad (2.62)$$

where $\mathcal{F}_1(\mathbf{f})$ and $\mathcal{F}_2(\mathbf{f})$ are lower semi-continuous, proper convex functions on $\mathbb{R}^{N'}$. The Fenchel conjugate of $\mathcal{F}(\mathbf{f})$ is defined as

$$\mathcal{F}^*(\mathbf{v}) = \sup_{\mathbf{f} \in \mathbb{R}^{N'}} \{ \langle \mathbf{v}, \mathbf{f} \rangle - \mathcal{F}(\mathbf{f}) \} \quad (2.63)$$

According to Fenchel-Young inequality:

$$\mathcal{F}(\mathbf{f}) + \mathcal{F}^*(\mathbf{v}) \geq \langle \mathbf{f}, \mathbf{v} \rangle \quad (2.64)$$

The equality holds if and only if $\mathbf{v} \in \partial \mathcal{F}(\mathbf{f})$

With the Fenchel conjugate of $\mathcal{F}_1$ and $\mathcal{F}_2$, the dual problem turns to be

$$\min_{\mathbf{v} \in \mathbb{R}^N} \mathcal{F}_2^*(\mathbf{v}) - \mathcal{F}_1^*(\mathbf{v}) \quad (2.65)$$

Combining Eq.2.62 and Eq.2.65, with affine minorizations, the solution to Eq.2.61 will be found iteratively via

$$\begin{aligned} \mathbf{v}^k &\in \partial \mathcal{F}_2(\mathbf{f}^k) \\ \mathbf{f}^k &\in \partial \mathcal{F}_1^*(\mathbf{v}^k) \end{aligned} \quad (2.66)$$

2.5 Deep learning introduction

2.5.1 Generalities

Recently, deep learning approaches have achieved great success in image processing with tasks such as classification (Krizhevsky et al., 2012), segmentation (Ronneberger et al., 2015), denoising (Zhang et al., 2017). They have been also proposed to solve inverse problems in imaging (McCann et al., 2017). In fact, both dictionary learning and deep learning belong to machine learning. However, deep learning approaches have two principal advantages which distinguish it from other approaches: powerful representation ability and much developed parallel calculation.

1. Deep learning learns high-level features of the data in a similar way to the visual processing in human brain using multiple layers of neurons with non linearities (LeCun et al., 2015). A large quantity of parameters is used with deep learning networks in order to reveal implicit information and to derive $\mathcal{A}^+$ which is an approximate inverse operator of $\mathcal{A}$. Yet, the efficiency of deep learning approach is related to the amount of data. A large amount of data could boost the performance of deep learning, on the contrary, a limited number of data limits its performance.
2. Parallel calculation has been much developed in deep learning framework, such as Tensorflow (Abadi et al., 2016), MXNet (Chen et al., 2015), Caffe (Jia et al., 2014). The user can directly benefit from high speed parallel computation without knowing GPU architecture and low-level GPU programming.

2.5.2 Network units

The deep learning network is a multi layers neuron network. Its first layer is named as input layer, the last layer is named as output layer. The intermediate layers are named as hidden layers. Fig.2.3 illustrates a classic 5 layers Multi Layer Perceptron (MLP) neuron network. Each neuron in the network consists of linear transformation followed by a point-wise nonlinear activation function.

Linear transformation

Every neuron at layer l is connected to all the neurons at layer $l + 1$ with weight $\theta_{i,j}^l$, i is the neuron index at layer l , j is the neuron index at the latter layer. If the output of

layer l is denoted as $\mathbf{L}^l$, the activation function at layer l is σ^l , the output of layer $l + 1$ as $\mathbf{L}^{l+1}$, then

$$\mathbf{L}^{l+1} = \sigma^l(\boldsymbol{\theta}^l \mathbf{L}^l) \quad (2.67)$$

The operation between $\boldsymbol{\theta}^l$ and $\mathbf{L}^l$ is matrix multiplication. In convolutional neural network (CNN), the linear operation is in form of convolution. MLP is one of the most classic deep learning networks. As a matter of fact, one layer network in MLP can be reduced as a 1×1 convolution layer in CNN for 2D image processing.

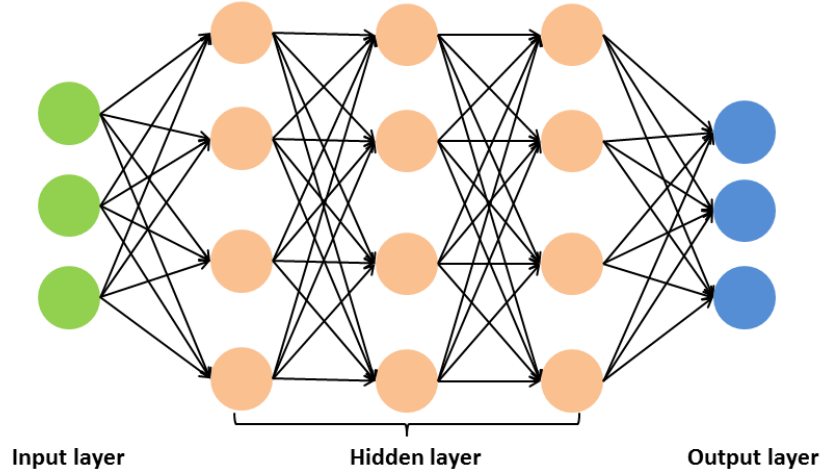


Figure 2.3: Multi Layer Perceptron (MLP) neuron network. In the training stage, input and output are fed with training data, the network is optimized by minimizing a user-selected function with respect to parameters within the network. Each solid circle represents a neuron which consists of linear transformation and nonlinear activation function.

Activation functions

Activation functions are essential for the network since they introduce nonlinear factors to the network. Here we present some activation functions which are used in the network. Rectified Linear Unit is illustrated in Fig.2.4(a). It linearly rectifies the nonnegative values and removed negative parts. PReLU (Parametric Rectified Linear Unit) is developed from ReLU activation, as shown in Fig.2.4(b). It rectifies the negative and nonnegative values in different ways.

PReLU is defined as

$$PReLU(x) = \begin{cases} x, & \text{if } x \geq 0 \\ ax, & \text{otherwise.} \end{cases}$$

where a is adapted during the training process.

The sigmoid function is used for 2-label classification tasks. Fig.2.4 displays its curve and it is defined as:

$$sigmoid(x) = \frac{1}{1 + e^{-x}} \quad (2.68)$$

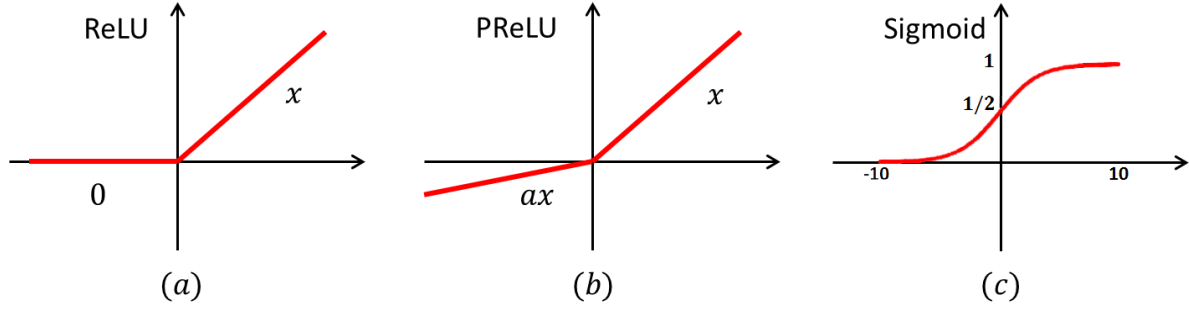


Figure 2.4: Illustration of three activation functions: ReLU, PReLU and sigmoid activation functions. The coefficient for the negative parts in ReLU is zero, whereas in PReLU, it is parametrized by a . Sigmoid function is used for two label classification.

It can be seen that when $x < -10$, the output of sigmoid function is close to 0, when $x > 10$, its output is close to 1.

2.5.3 CNN

Among the deep learning networks, many architectures have been proposed: MLP, basic CNN, recurrent neural networks, recursive network, stacked denoising auto-encoders, generative adversarial networks and so forth (LeCun et al., 2015).

CNN is a workhorse of deep learning, particularly of image processing. In this work, we focus on CNN. With CNN, the inverse operator of $\mathcal{A}$ is approximated by a sequence of filtering operations alternating with nonlinear operations. The optimization is performed on the parameters in the network, such as the weights of the filters, parameters involved in the activation functions, and so forth.

The convolution operation allows network to detect the same feature in different regions of image. It seems that filters at one CNN layer detect local features of image, while multi-layer CNN allows to increase the perception field and to synthesize the features abstracted at previous layers. Moreover, CNN reduces the number of weights by sharing them between network's neurons, which results in a considerable memory reduction.

2.5.4 Optimization of the network

The application of deep learning in super resolution has been broadly discussed in the literature (Dong et al., 2014, Kim et al., 2016a, Kim et al., 2016b, Ledig et al., 2016, Shi et al., 2016, Lai et al., 2017, Zhang et al., 2018, Lim et al., 2017, Tai et al., 2017). In order to clarify the optimization of the network, we take super resolution as an example to introduce some basic strategies.

Let $\mathbf{g}$ denote a low-resolution image and $\mathbf{f}^*$ a high-resolution image. Given a training dataset $(\mathbf{g}_i, \mathbf{f}_i^*)$, our goal is to learn $\mathcal{A}^+$ which predicts values $\mathbf{f} = \mathcal{A}^+ \mathbf{g}$ so that $\mathbf{f} \approx \mathbf{f}^*$. $\mathcal{F}$ is a user-selected loss function used to optimize the parameters in the network, and it varies with task types. Moreover, it's possible to add regularization term in $\mathcal{F}$, for the purpose of reducing the over-fitting risk or injecting priors. Generally speaking, a

parametric approximate inverse $\mathcal{A}_\theta^+ : \mathcal{Y} \rightarrow \mathcal{X}$ is learned by solving:

$$\arg \min_{\theta \in \Theta} \sum \mathcal{F}(f_i^*, \mathcal{A}_\theta^+(g_i)) + \mathcal{G}(\theta) \quad (2.69)$$

where Θ is the set of possible parameters. A function $\mathcal{G}$ may regularize the problem. The function $\mathcal{F}$ is a loss function measuring the error. The value of the function $\mathcal{F}$ is computed via forward propagation, and the gradients of $\mathcal{F}$ with respect to parameters in the network are determined via back-propagation (Rumelhart et al., 1986). Stochastic gradient descent method accelerates back propagation by processing data in small batch (LeCun et al., 2015). The parameters are optimized by successive forward and backward pass in the network.

The fitting capacity is described by the 'bias', which is the expectation of error on the training set. The generalization capacity is evaluated by the 'variance', which is the evaluated on the test set. With the increase of the network size, the bias tends to decrease while the variance tends to increase, which can be interpreted as the network evolves from underfitting to overfitting. It is generally admitted that, in a network, a trade off must be found between the bias and the variance in order to prevent overfitting along with a small generalization error.

Many factors influence the performance of the network: the network architecture, the training set, the optimization procedure. The architecture of the network varies with the type of task. A large dataset in which the data follows the same distribution boosts the performance of the network. A good optimization procedure can increase the efficiency of the training process, or improve the accuracy of the estimation.

In this section, a brief introduction of deep learning is given. We hope it gives a general map to readers about the deep learning. More details can be found in chapter 6 and chapter 7.

2.6 Discussion

In this chapter, we have defined ill-posed inverse problems and the state of art for super resolution methods in which we will be interested. We considered four classes of methods to solve it: spatial interpolation, super resolution in frequency domain, regularization methods, machine learning methods such as dictionary learning and deep learning. Moreover, some algorithms solving convex and nonconvex problems have been presented. We may note that there are strong links between the different approaches or concepts, for instance, machine learning methods use some principles of regularization schemes and methods for convex and non convex optimization.

Chapter 3

Determination of the kernel

3.1 Introduction

The super resolution problem that we want to solve in our work includes both the effect of a convolution kernel and a subsampling operator. Deconvolution is a major area in inverse problems in imaging and has been extensively studied. Many methods and algorithms have been proposed to solve it. However, it is also known that the results are very sensitive to the choice of the kernel. It is thus crucial to have an estimation of the blurring kernel as good as possible.

In the deconvolution problems, the origins of blur are diverse and various methods have been developed to estimate blur kernels (Reeves and Mersereau, 1992, Fergus et al., 2006, Joshi et al., 2008). The determination of the kernel is an ill-posed problem and the methods differs by the data term and the regularization term used, and also by the general framework, deterministic or bayesian. Some assumptions have to be made about the type of noise corrupting the images. In the literature, if the kernels are assumed to follow a certain distribution, such as Gaussian or Laplacian, the methods are classified as semi-blind deconvolution (Reeves and Mersereau, 1992, Fergus et al., 2006, Peyrin et al., 2015). If the kernels are not parametrized, the methods are classified as blind deconvolution (Joshi et al., 2008, Kotera et al., 2013, Xu et al., 2014, Hanocka and Kiryati, 2015, Zhao et al., 2015, Tofighi et al., 2016, Xu et al., 2014). Common blind deconvolution methods first determine the kernel, then apply the kernel to recover images. Some researchers propose to iteratively optimize the kernel and recover image (Hanocka and Kiryati, 2015, Zhao et al., 2015).

J.Reeves *et al.* in (Reeves and Mersereau, 1992) used the cross validation to determine the parameters of the blur. Fergus *et al.* in (Fergus et al., 2006) attempted to use Bayesian approach to deblur single images degraded by camera shaking. In this work, the authors assume that the degradation kernel is a linear combination of exponential distributions, which is a way to add a priori information so that the inverse problem less ill-posed.

Joshi *et al.* in (Joshi et al., 2008) proposed to first predict the latent images by sharpening the edges of the blurred images, then these authors determine the degradation kernel and the input image. The final deconvolution images are obtained with the resolved kernels via Lucy-Richardson algorithm. Moreover, the blur determination can be regularized with a finite dimensional approximation (Denis et al., 2015) or patch-based priors (Sun et al., 2013). Some specific methods for motion blurred images have been

studied in (Cannon, 1976, Chang et al., 1991).

In the literature, blind deconvolution has been used to remove the effects of motion or rotation or other degradation factors. Blind deconvolution determines the kernel based on the input images (low resolution), no matter it is degraded by motion or other factors. In our problem, we previously considered a blind TV approach based on a Gaussian prior of the Point Spread Function (PSF). However, the experimental images were not successfully recovered since the PSF is more complex than expected. For this reason, in this chapter, we attempt to derive the kernel based on known high and low resolution images.

We propose to use a regularization functional which is based on Mutual Information to estimate the kernel and to get the optimal solution with gradient descent approach (Li et al., 2017b). The kernels obtained with this approach and L_2 norm optimization are then compared to solve a super resolution problem on real data with the TV regularization. Their efficiency to improve the *in vivo* HR-pQCT image quality is evaluated with respect to the accuracy of the quantitative parameters of bone micro-structure.

This chapter is organized as follows. In Sect.3.2, we present the image observation model and the regularization functional used for the kernel estimation. The optimization methods are detailed as well. In Sect.3.3, we present the application of the TV regularization in super resolution problem. Afterwards, apply the resolved kernel in the super resolution problem with TV regularization approach. Numerical results are presented in the Sect.3.4. Concluding remarks are given at the end of this chapter.

3.2 The inverse problem of the blurring kernel estimation

Similar to the definition in 2.1, $\mathbf{g} \in \mathbb{R}^N$ is denoting the 3D low resolution image obtained from HR-pQCT, $\mathbf{f}^* \in \mathbb{R}^{N'}$ the 3D ground truth imaged from μ CT, the isotropic under-sampling factor is 2, thus $N' = 8N$. We define the blurring kernel as $\mathcal{K} : \mathbb{R}^{N'} \rightarrow \mathbb{R}^{N'}$, and the undersampling operator as $\mathcal{U} : \mathbb{R}^{N'} \rightarrow \mathbb{R}^N$. The forward observation model is formulated as

$$\begin{aligned} \mathbf{g} &= \mathcal{A}\mathbf{f}^* + \mathbf{n} \\ &= \mathcal{U}(\mathcal{K} * \mathbf{f}^*) + \mathbf{n} \end{aligned} \quad (3.1)$$

"*" is the convolution operator, $\mathbf{n}$ is the additive noise. The undersampling operator $\mathcal{U}$ is illustrated in 3.1:

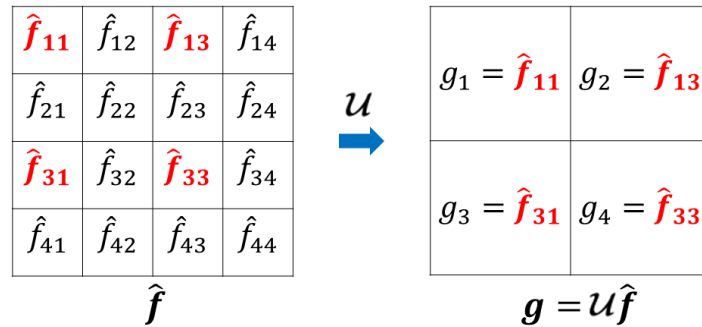


Figure 3.1: Demonstration of the undersampling operator $\mathcal{U}$ with factor equal to 2.

Since the intensities of $\mathbf{f}^*$ and $\mathbf{g}$ are not in the same ranges, we introduce a matrix $\mathbf{c} = c \times \mathbb{1}$ to adapt this discrepancy:

$$\begin{aligned}\mathbf{g} &= \mathcal{A}\mathbf{f}^* + \mathbf{n} \\ &= \mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c})) + \mathbf{n}\end{aligned}\tag{3.2}$$

where $c \in \mathbb{R}$ is updated at each iteration, $\mathbb{1} \in \mathbb{R}^{N'}$ is a matrix whose values are all equal to 1.

In the following, the subsection 3.2.1 describes the determination of the kernel via L_2 norm minimization, the subsection 3.2.2 presents how to determine the kernel based on the mutual information notion.

3.2.1 L_2 norm minimization

The determination of the blurring kernel $\mathcal{K}$ from a couple of high and low resolution images can be considered as a linear inverse problem. The usual approach can be formulated as:

$$\arg \min_{\mathcal{K}, \mathbf{c}} \mathcal{J}_1(\mathcal{K}, \mathbf{c}) = \|\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c})) - \mathbf{g}\|_2^2 + \alpha \mathcal{R}(\mathcal{K})\tag{3.3}$$

where α is a regularization parameter and $\mathcal{R}$ a regularization functional (Scherzer et al., 2009). The regularization term $\mathcal{R}(\mathcal{K}) = \|\nabla \mathcal{K}\|_2^2$ is chosen to penalize the fast variations of the kernel.

The search for the minimizer of the regularization functionals is implemented with a gradient descent strategy. For both regularization functionals we use an alternate minimization method with respect to $\mathcal{K}$ and $\mathbf{c}$. The gradients of the functional $\mathcal{J}_1$ with respect to $\mathcal{K}$ and $\mathbf{c}$ are:

$$\begin{aligned}\frac{\partial \mathcal{J}_1(\mathbf{c}, \mathcal{K})}{\partial \mathcal{K}} &= (\mathbf{f}^* + \mathbf{c}) * \mathcal{U}((\mathbf{f}^* + \mathbf{c}) * \mathcal{K}) - (\mathbf{f}^* + \mathbf{c}) * \mathbf{g} + \alpha \Delta \mathcal{K} \\ \frac{\partial \mathcal{J}_1(\mathbf{c}, \mathcal{K})}{\partial \mathbf{c}} &= 2c \|\mathcal{U}(\mathbb{1} * \mathcal{K})\|^2 + 2\langle \mathcal{U}(\mathcal{K} * \mathbf{f}^*) - \mathbf{g}, \mathcal{U}(\mathbb{1} * \mathcal{K}) \rangle\end{aligned}$$

where $\langle \cdot, \cdot \rangle$ is the inner product, the matrix $\mathbb{1}$ has the same size as the kernel $\mathcal{K}$, the spatial convolution $\mathbb{1} * \mathcal{K}$ is calculated in frequency domain.

3.2.2 Mutual Information based approach

Why Mutual Information?

Other statistical dissimilarity measures (Hermosillo et al., 2002) may be more efficient than the L_2 norm for enhancing the similarity between the histogram distribution of $\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))$ and $\mathbf{g}$. To this aim, we propose to exploit the mutual information (MI) between $\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))$ and $\mathbf{g}$. The estimated kernel $\mathcal{K}$ will then be obtained with the minimization of the regularization functional $\mathcal{J}_2(\mathcal{K}, \mathbf{c})$:

$$\mathcal{J}_2(\mathcal{K}, \mathbf{c}) = -\text{MI}(\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c})), \mathbf{g}) + \alpha \|\nabla \mathcal{K}\|_2^2\tag{3.4}$$

where $\text{MI}(\cdot, \cdot)$ calculates the mutual information of two input variables. The MI between $\mathcal{U}(\mathbf{f} * \mathcal{K})$ and $\mathbf{g}$ is based on their joint probability function. It gives an estimate

of their statistical dependence and has been used for multi-modal image matching (Hermosillo et al., 2002). It is defined by:

$$\text{MI}(\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c})), \mathbf{g}) = \sum_{i,j} P_{conj}(i, j) \cdot \ln \frac{P_{conj}(i, j)}{P_{\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))}(i) P_{\mathbf{g}}(j)} \quad (3.5)$$

where $P_{conj}(i, j)$ is the conjugate probability distribution of $\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))$ and $\mathbf{g}$ for intensity values i and j , and $P_{\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))}(i)$, $P_{\mathbf{g}}(j)$ are the marginal probability distribution of $\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))$ and $\mathbf{g}$. In order to evaluate the MI, we consider a Parzen estimator (Parzen, 1962, Bosq, 2012) for the joint probability function with a Gaussian of variance β , and a normalization constant $|\Omega| = \sum_{i,j} P_{conj}(i, j)$:

$$P_{conj}(i, j) = \frac{1}{|\Omega|} \int_{\Omega} \exp\left(-\frac{(\mathcal{U}(\mathcal{K} * (\mathbf{f}^*(x) + \mathbf{c})) - i)^2}{\beta^2}\right) \cdot \exp\left(-\frac{(\mathbf{g}(x) - j)^2}{\beta^2}\right) dx$$

The marginal probability distributions can be obtained as:

$$\begin{aligned} P_{\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))}(i) &= \sum_j P_{conj}(i, j) \\ P_{\mathbf{g}}(j) &= \sum_i P_{conj}(i, j) \end{aligned}$$

Minimization algorithm

Similarly to the L_2 norm minimization method, the minimization of $\mathcal{J}_2$ is implemented with a gradient descent strategy and an alternate minimization.

The **MI** is a nonlinear functional of the kernel $\mathcal{K}$. Given an initial estimate $\mathcal{K}_0$, denote **NMI**=-**MI**, the optimization of the kernel is equivalent to the solution of the initial value problem:

$$\begin{cases} \frac{d\mathcal{K}}{dt} = -\frac{\partial \mathcal{J}_2(\mathcal{K}, \mathbf{c})}{\partial \mathcal{K}} \\ \quad = -\left(\frac{\partial \text{NMI}(\mathcal{K}, \mathbf{c})}{\partial \mathcal{K}} - \alpha \Delta \mathcal{K}\right) \\ \mathcal{K}(0) = \mathcal{K}_0 \end{cases} \quad (3.6)$$

where $\frac{\partial \mathcal{J}_2(\mathcal{K}, \mathbf{c})}{\partial \mathcal{K}}$ is the gradient of $\mathcal{J}_2$ with respect to $\mathcal{K}$. We will admit in this work that this first order differential equation is well-defined.

The first variation of the **MI** at $\mathcal{K}$ in the direction of $\mathbf{k}$ is defined by:

$$\delta_{\mathbf{k}} \text{NMI}(\mathcal{K}) = \left. \frac{\partial \text{NMI}(\mathcal{K} + \epsilon \mathbf{k})}{\partial \epsilon} \right|_{\epsilon=0} \quad (3.7)$$

and the gradient is obtained with:

$$\delta_{\mathbf{k}} \text{NMI}(\mathcal{K}) = \left\langle \frac{\partial \text{NMI}(\mathcal{K}, \mathbf{c})}{\partial \mathcal{K}}, \mathbf{k} \right\rangle \quad (3.8)$$

Considering the fact that P_{conj} is a probability distribution, we have (Hermosillo et al., 2002):

$$\delta_{\mathbf{k}} \text{NMI}(\mathbf{k}) = - \sum_{i,j} \delta_{\mathbf{k}} P_{conj}(i, j) \ln \frac{P_{conj}(i, j)}{P_{\mathcal{U}(\mathcal{K} * (\mathbf{f}^* + \mathbf{c}))}(i) P_{\mathbf{g}}(j)} \quad (3.9)$$

where

$$\delta_{\mathbf{k}} P_{conj} = -\frac{2}{\beta^2} \int_{\Omega} \{ \mathbf{u}(\mathbf{K} * (\mathbf{f}^*(x) + \mathbf{c})) (\mathbf{u}(\mathbf{K} * (\mathbf{f}^*(x) + \mathbf{c})) - i) \exp(-\frac{(\mathbf{u}(\mathbf{K} * (\mathbf{f}^*(x) + \mathbf{c})) - i)^2 + (\mathbf{g}(x) - j)^2}{\beta^2}) \} dx \quad (3.10)$$

Based on this formula, we can show that the gradient of the **NMI** (negative **MI**) with respect to $\mathbf{K}$ is given by:

$$\begin{aligned} \frac{\partial \text{NMI}(\mathbf{K}, \mathbf{c})}{\partial \mathbf{K}} = & 2(\mathbf{f}^* + \mathbf{c}) * \sum_{i,j} \{ (\mathbf{u}(\mathbf{K} * (\mathbf{f}^*(x) + \mathbf{c})) - i) \cdot \\ & \exp(-\frac{(\mathbf{u}(\mathbf{K} * (\mathbf{f}^*(x) + \mathbf{c})) - i)^2 + (\mathbf{g}(x) - j)^2}{\beta^2}) \cdot \\ & \ln \frac{P_{conj}(i, j)}{P_{\mathbf{u}(\mathbf{K} * (\mathbf{f}^* + \mathbf{c}))}(i) P_{\mathbf{g}}(j)} \} \end{aligned} \quad (3.11)$$

On the other hand, the gradient of the negative **MI** with respect to the constant $\mathbf{c}$ is given by:

$$\frac{\partial \text{NMI}(\mathbf{K}, \mathbf{c})}{\partial \mathbf{c}} = \frac{\int \int \int \mathcal{T}(i, j, x) \frac{2\mathbf{u}(\mathbf{K} * \mathbb{1}) \cdot (\mathbf{u}(\mathbf{K} * (\mathbf{f}^* + \mathbf{c})) - i)}{-\beta^2} didjdx}{\int \int \int \mathcal{T}(i, j, x) \frac{(\mathbf{u}(\mathbf{K} * \mathbb{1}))^2}{-\beta^2} didjdx} \quad (3.12)$$

x denotes the intensity of one voxel on the image, and

$$\mathcal{T}(i, j, x) = \exp(-\frac{(\mathbf{u}(\mathbf{K} * (\mathbf{f}^* + \mathbf{c})) - i)^2 + (\mathbf{g} - j)^2}{\beta^2}) \cdot \ln \frac{P_{conj}(i, j)}{P_{\mathbf{u}(\mathbf{K} * (\mathbf{f}^* + \mathbf{c}))}(i) P_{\mathbf{g}}(j)} \quad (3.13)$$

With Eq.3.11 and Eq.3.12, μ_1 , μ_2 and α are three scalar parameters, $\mathbf{K}$ and $\mathbf{c}$ are updated by

$$\begin{cases} \mathbf{K}^{k+1} = \mathbf{K}^k - \mu_1 \left(\frac{\partial \text{NMI}(\mathbf{K}^k, \mathbf{c}^k)}{\partial \mathbf{K}^k} - \alpha \Delta \mathbf{K}^k \right) \\ \mathbf{c}^{k+1} = \mathbf{c}^k - \mu_2 \frac{\partial \text{NMI}(\mathbf{K}^{k+1}, \mathbf{c}^k)}{\partial \mathbf{c}^k} \end{cases} \quad (3.14)$$

In this section, we have introduces two ways to estimate kernels: one is based on L_2 norm minimization, the other is based on mutual information optimization. For the purpose of comparing the two resolved kernels, we're going to apply them in TV regularization to solve super resolution problems.

3.3 Super resolution with the TV regularization

With the same forward problem model described in Eq.3.2, the objective in the super resolution problem is to find a solution $\mathbf{f}$ which is close to the ground truth $\mathbf{f}^*$.

The TV regularization is based on the gradient information and tends to give a smooth solution. $\mathbf{D}_i$ calculates the gradient of image $\mathbf{f}$ as voxel i . The TV regularization is defined as

$$\mathcal{R}(\mathbf{f}) = \sum_i \|\mathbf{D}_i \mathbf{f}\|_2 \quad (3.15)$$

After the combination of the data fitting term and the TV regularization, the objective functional is written as

$$\mathcal{J}(\mathbf{f}) = \frac{\mu}{2} \|\mathcal{A}\mathbf{f} - \mathbf{g}\|_2^2 + \sum_i \|\mathcal{D}_i \mathbf{f}\|_2 \quad (3.16)$$

where μ is a regularization parameter balancing the effects of the data fitting term and the regularization term. When μ is very small, the solution is over smoothed and many original information are lost. On the contrary, if μ is very huge, no smoothing effect can be observed, thus the image is quite noisy.

A super resolution image can be obtained by minimizing the functional of Eq.3.16. Despite that the TV regularization is a convex but nonsmooth term, Eq.3.16 can be optimized by ADMM (in Sect.2.3.2). We consider the augmented Lagrangian functional

$$\mathcal{L}_A(\mathbf{f}, \mathbf{u}, \boldsymbol{\lambda}) = \frac{\mu}{2} \|\mathcal{A}\mathbf{f} - \mathbf{g}\|_2^2 + \sum_i \{ \|\mathbf{u}_i\|_2 - \boldsymbol{\lambda}_i^t (\mathbf{u}_i - \mathcal{D}_i \mathbf{f}) + \frac{\beta}{2} \|\mathbf{u}_i - \mathcal{D}_i \mathbf{f}\|_2^2 \} \quad (3.17)$$

$$\text{s.t. } \forall i, \mathcal{D}_i \mathbf{f} = \mathbf{u}_i$$

where $\mathbf{u}_i$ is the voxel of $\mathbf{u}$ at the index i . $\boldsymbol{\lambda}_i$ is the Lagrangian multiplier, $\beta \in \mathbb{R}^+$ is the Lagrangian parameter. The saddle point of the $\mathcal{L}_A$ is estimated with successive updates.

1. Update of $\mathbf{f}^{k+1}$:

$$(\mu \mathcal{A}^t \mathcal{A} + \beta \sum_i \mathcal{D}_i^t \mathcal{D}_i) \mathbf{f}^{k+1} = \mu \mathcal{A}^t \mathbf{g} + \sum_i \mathcal{D}_i^t (\beta \mathbf{u}_i^k - \boldsymbol{\lambda}_i^k) \quad (3.18)$$

2. Update of $\mathbf{u}_i^{k+1}$:

$$\begin{aligned} \mathbf{u}_i^{k+1} &= \mathcal{S}_\beta(\mathcal{D}_i \mathbf{f}^{k+1} + \frac{\boldsymbol{\lambda}_i^k}{\beta}) \\ \mathcal{S}_\beta(\mathbf{x}) &= (\max(1 - \frac{1}{\beta \|\mathbf{x}\|_2 + \epsilon}, 0)) \mathbf{x} \end{aligned} \quad (3.19)$$

ϵ is a small number ensuring the stability of the algorithm in case of $\|\mathbf{x}\|_2 = 0$.

3. Update of $\boldsymbol{\lambda}_i^{k+1}$:

$$\boldsymbol{\lambda}_i^{k+1} = \boldsymbol{\lambda}_i^k - \beta(\mathbf{u}_i^{k+1} - \mathcal{D}_i \mathbf{f}^{k+1}) \quad (3.20)$$

Iterations are stopped when

$$\frac{\|\mathbf{f}^{n+1} - \mathbf{f}^n\|_2}{\|\mathbf{f}^{n+1}\|_2} \leq 10^{-3} \quad (3.21)$$

3.4 Numerical experiments

3.4.1 Simulation details

The tests for the determination of the blurring kernel were performed on 3D experimental images of trabecular bone obtained from HR-pQCT (Boutroy et al., 2005) and μ -CT (Burghardt et al., 2011). The voxel sizes of the HR-pQCT and μ CT images are $82\mu\text{m}$

and $41\mu\text{m}$ respectively. The tested high and low resolution images were respectively cropped at $N' = 200 \times 200 \times 200$ voxels and $N = 100 \times 100 \times 100$ voxels.

The initial kernel $\mathcal{K}_0$ chosen is a Gaussian with a standard deviation of 4. Several support sizes for the kernel have been investigated with odd number ranging from 5 - 13.

In TV super resolution restoration, parameters are chosen to obtain the best decrease of the regularization functional, $\mu = 10$, $\beta = 7$ for both kernels.

3.4.2 Criteria for image quality

To quantify the quality of the super resolved images, we have used different metrics. The PSNR is a classical criterion to evaluate the similarity between gray level images. Since the final goal of the study is to assess bone structure, we used three other measures calculated on the binarized bone images: the DICE index, BV/TV and the connectivity density (Conn.D)(Odgaard, 1997).

Denoting $\mathbf{f}^*$ the ground truth image, $\mathbf{f}$ the reconstructed image, $\mathbf{f}_b^*$ and $\mathbf{f}_b$ are their corresponding binary images, the criteria mentioned above are defined as:

$$\begin{aligned} \text{MSE} &= \frac{1}{N'} \sum_{i=1}^{N'} (\mathbf{f}_i^* - \mathbf{f}_i)^2 \\ \text{PSNR} &= 10 \log_{10} \frac{(\max(\mathbf{f}^*) - \min(\mathbf{f}^*))^2}{\text{MSE}} \\ \text{DICE} &= \frac{2|\mathbf{f}_b \cap \mathbf{f}_b^*|}{|\mathbf{f}_b| + |\mathbf{f}_b^*|} \\ \text{BV/TV} &= \frac{\text{bone volume}}{\text{total volume}} \\ \text{Conn.D} &= \frac{\beta_1}{\text{total volume}} = \frac{\beta_0 + \beta_2 - \varepsilon}{\text{total volume}} \end{aligned} \tag{3.22}$$

where in DICE index, $|\cdot|$ counts the number of 1. Remind that β_1 is the first order Betti number. It is defined as the maximal number of cuts that can be applied to a domain so that it remains connected. For instance, β_1 for a line is 0 and β_1 for a torus is 2. β_0 and β_2 are the number of cavities and connected components in the image, and ε is the Euler number(Toriwaki and Yonekura, 2002).

The DICE index is a typical metric to assess the similarity of binary images. The BV/TV and Conn.D are two common parameters in bone studies: BV/TV reflects the bone volume fraction and the Conn.D represents topological characteristics of bone micro-architecture. The binarization is performed by the Otsu (Dice, 1945) threshold which is determined based on image histogram.

All the parameters defined in Eq.3.22 are used during the whole thesis to evaluate different reconstruction approaches.

3.4.3 Numerical results

The super resolved images obtained by TV regularization with these two optimized kernels are displayed in Fig.3.2 together with the low and high resolution images. Fig.3.2(c)(d) display more structural details compared with the low resolution image (b), but the reconstructed structures are thick compared with the high resolution image (a).

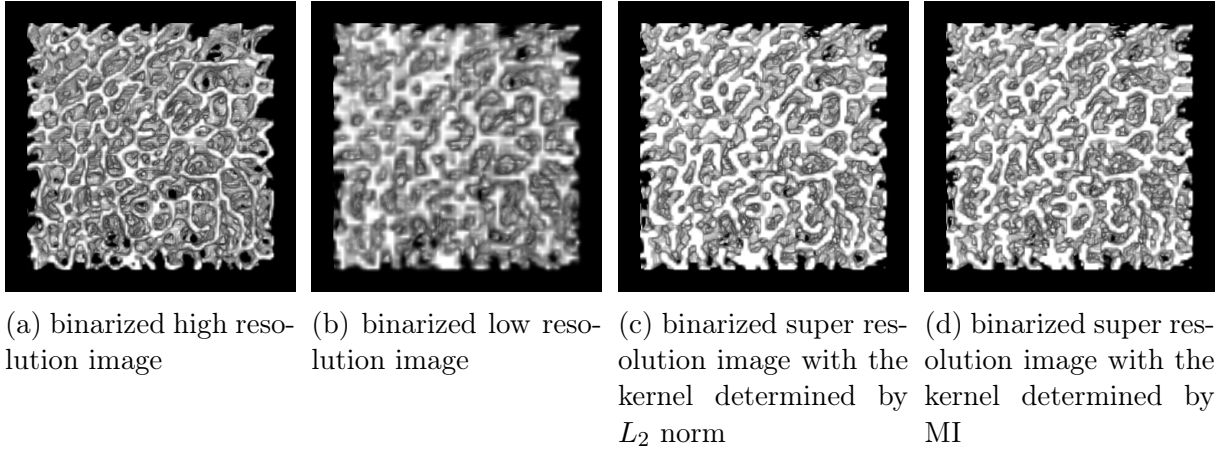


Figure 3.2: illustration of binarized high resolution, low resolution, L_2 kernel restored image and MI kernel restored image.

The evolutions of the PSNR, DICE, BV/TV, Conn.D as well as the data fitting term as a function of the number of iterations are illustrated in Fig.3.3-3.7. The PSNR reflects the similarity of two images in gray level, while DICE, BV/TV and connectivity density compare binary images. As expected, we obtain a decrease of the reconstruction errors evaluated with the PSNR or the DICE with the two methods. The kernel obtained with L_2 norm minimization performs slightly better in term of the PSNR, whereas the one derived from MI is better regarding the DICE. Generally speaking, the performance of both kernels are comparable.

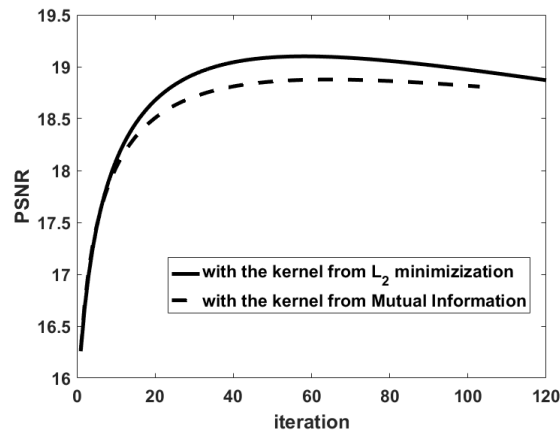


Figure 3.3: Evolution of PSNR as a function of the number of iteration for the two kernels.

The obtained super-resolution images have thinner structures than the low resolution images, but the reconstructed structures still have a larger bone mass than high resolution images. As shown in Fig.3.5, both methods improves image quality with respect to BV/TV, but this ratio is still too high in the reconstructed images. Conn.D is a topological criterion. Even though it has been improved as displayed in Fig.3.6, the final divergence tail shows that the determined kernels did not appropriately recover the degraded topological structures. In addition, the kernel from L_2 minimization changes the topological

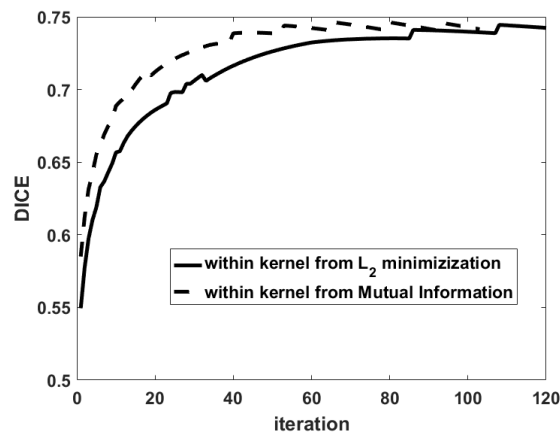


Figure 3.4: Evolution of DICE as a function of the number of iteration for the two kernels.

structure of the super-resolution image more slowly than the one from MI optimization. Fig. 3.7 illustrates the data term evolution. This figure also shows the convergence of the TV reconstructions with the two optimized kernels. In general, the kernel deduced from L_2 norm is more efficient than MI considering Conn.D and local artifacts, whereas the kernel from MI is slightly better regarding BV/TV and data fitting term.

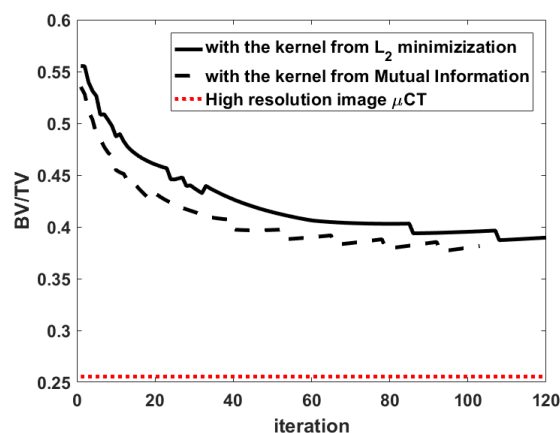


Figure 3.5: Evolution of BV/TV as a function of the number of iteration for the two kernels.

3.5 Conclusion

The determination of the kernel is a major issue in image restoration. In this chapter, we have compared two methods of determination of the blurring kernel based on couples of experimental low and high resolution images. Similar results were obtained with both kernels derived by minimizing a L_2 norm or MI. The kernel from L_2 norm preserves a better connectivity density, and the one from MI has a better performance on the BV/TV. One possible improvement for our MI gradient descent method is to consider local information (Hermosillo et al., 2002). Hermosillo *et al.* proposed a local MI method,

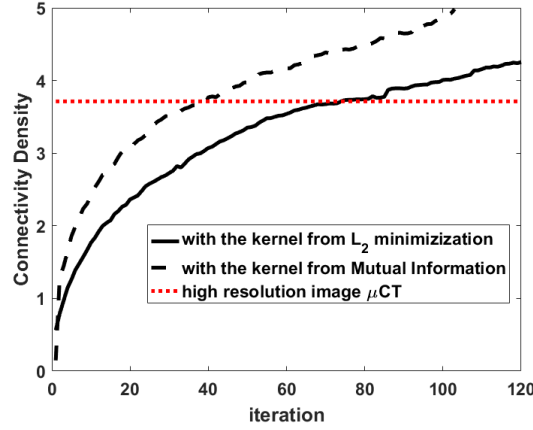


Figure 3.6: Evolution of connectivity density as a function of the number of iteration for the two kernels.

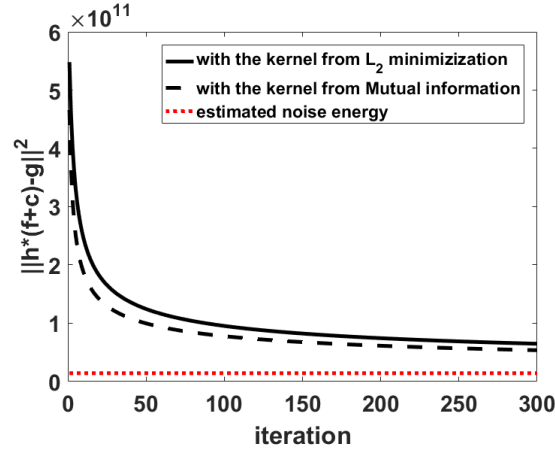


Figure 3.7: Evolution of data fitting term as a function of iteration. In order to show the whole evolution of the data fitting term, this figure didn't stop iteration by relative changes, which is $\|\mathbf{f}^{n+1} - \mathbf{f}^n\|_2 / \|\mathbf{f}^{n+1}\|_2 \leq 10^{-3}$.

which will generate a spatially variant kernel that may improve the results. These results need to be further confirmed on other ROIs and other HR-pQCT images.

Further more, according to the display of TV super resolution images, the deconvolution approaches improve the HR-pQCT images quality, but it's still far away from our expectations. We consider to explore more priors to solve this super resolution task. In the next chapter, we integrate a new prior on the super resolution image in order to enhance the contrast of the super resolution images.

Chapter 4

Nonconvex and nonsmooth optimization to enhance the contrast of images

4.1 Introduction

An image with a good contrast favors the segmentation task, particularly for the extraction of the trabecular micro-architecture from gray level CT images. The ground truth images provided by μ -CT have a quasi-binary nature and their histograms follow bimodal distributions. However, this characteristic is lost in low resolution images generated from HR-pQCT. Fig.4.1 illustrates the histogram of high and low resolution images.

We attempt to recover the bimodality of super resolution histograms, and an interesting approach is to include a double-well potential in the regularization term despite its nonconvexity. Based on the previous work (Li et al., 2017a), we propose to combine a double-well nonconvex constraint with TV nonsmooth regularization functional to enhance the bimodality of the resulting images.

Minimizing such a nonconvex and nonsmooth functional is difficult but numerical methods have been studied in recent years. In the second chapter 2.4, different basic frameworks for nonconvex nonsmooth optimization have introduced. Nonconvex regularizers have been investigated in (Foucart and Lai, 2009, Samson et al., 2000). Proximal operators (Li and Lin, 2015, Ochs et al., 2014) have been recently used to find critical points for nonconvex and non differentiable functionals. The ADMM has also been extended to nonconvex problems (Artina et al., 2013). Another possibility for the double-well potential is to solve the problem in Sobolev spaces (Bertozzi et al., 2007b, Bertozzi et al., 2007a). However, these methods have so far not been applied to the minimization problem considered here.

In this chapter, we propose three schemes to minimize the regularization functional combining TV and a double-well potential, then evaluate them for our problem. The first two approaches are based on a linearly constrained nonconvex and nonsmooth minimization method (Artina et al., 2013), and an accelerated gradient descent with proximal operators (Li and Lin, 2015). The third one is a new approach combining the Sobolev norms used in (Bertozzi et al., 2007b) with the TV regularization term.

This chapter is organized as follows. Sect.4.2 formulates the mixed TV and double-well

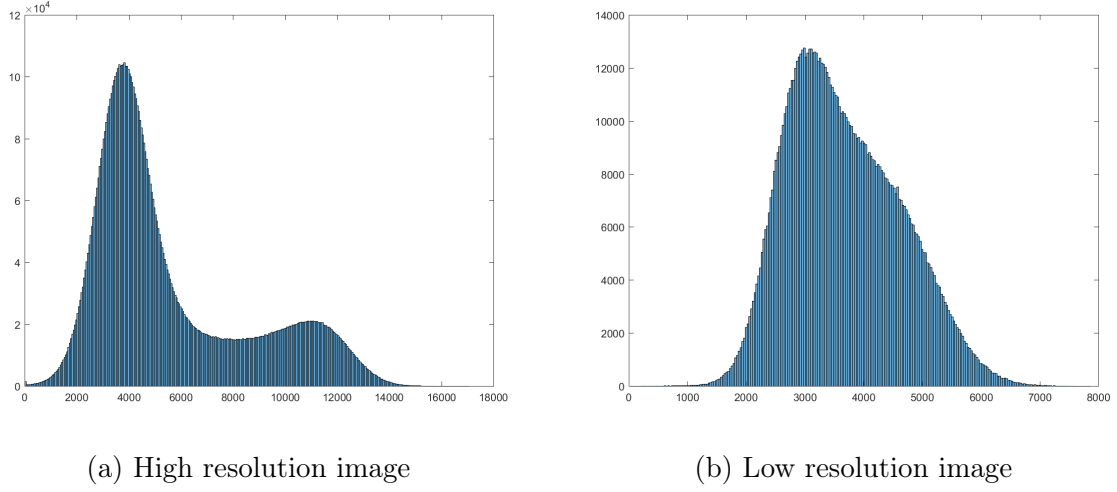


Figure 4.1: Typical histogram of high resolution (voxel size $41\mu\text{m}$) and low resolution (voxel size $82\mu\text{m}$) CT images. According to (a), s_1 and s_2 corresponds to the location of two peaks in the histogram of the high resolution image ($s_1 = 3750$, $s_2 = 11075$).

regularization functional. The three different approaches to minimize the regularization functional are detailed in Sect.4.3. Sect.4.4 presents results on experimental bone 3D HR-pQCT images. Concluding remarks will be given at the end of this chapter.

4.2 The inverse problem

With the same forward observation model than in Eq.2.2, the regularized functional to solve the super resolution is written as:

$$\mathcal{J}(f) = \|\mathcal{A}f - g\|_2^2 + \mu\mathcal{R}(f) \quad (4.1)$$

In chapter 3, $\mathcal{R}(\cdot)$ is the TV regularization. In this chapter, it includes TV regularization and the double-well potential. The double-well potential is drawn in Fig.4.2 and is written as

$$\mathcal{W}(f_i) = (f_i - s_1)^2(f_i - s_2)^2 \quad (4.2)$$

where $s_1, s_2 \in \mathbb{R}^+$ are the positions of the peaks of the bimodal histogram, and i is the index of the voxel location. The final objective functional is formulated as:

$$\mathcal{J}(f) = \|\mathcal{A}f - g\|_2^2 + \mu_1 \sum_i \|\mathcal{D}_i f\|_2 + \mu_2 \sum_i \mathcal{W}(f_i) \quad (4.3)$$

where μ_1, μ_2 are regularization parameters, $\mathcal{D}_i$ denotes the 3 dimensional gradients at voxel i .

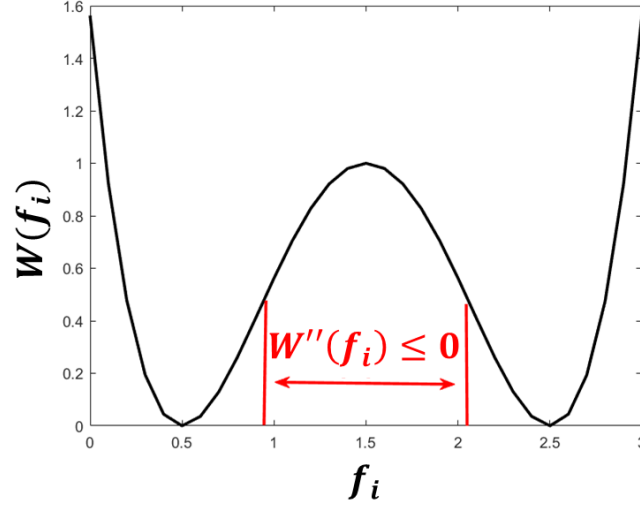


Figure 4.2: doubled-well potential function: $\mathcal{W}(f_i) = (f_i - 0.5)^2(f_i - 2.5)^2$

4.3 Algorithms for nonconvex and nonsmooth optimization

The first method, Linearly Constrained Nonconvex Nonsmooth Minimization (LCNNM), is based on a generalization of ADMM to nonconvex and nonsmooth problems (Artina et al., 2013). The second one is based on the Accelerated Proximal Gradient (APG) (Li and Lin, 2015). The last one is a new approach, it combines Cahn Hilliard model with TV (CH+TV) and is solved via a combination of the ADMM scheme with the minimization method of the Cahn-Hilliard potential based on the work presented in (Bertozzi et al., 2007b, Bertozzi et al., 2007a). The three approaches are not expected to converge to the same solutions even with the same initialization because the regularization functionals and the applied minimization strategies are distinct.

4.3.1 Linearly Constrained Nonsmooth and Nonconvex Minimization method (LCNNM)

The principle of the LCNNM method is to modify the nonconvex double-well potential to a convex functional by adding a quadratic term which is updated during the minimization (Artina et al., 2013).

We first consider the following functional:

$$\mathcal{J}(\mathbf{f}|\omega, \mathbf{v}) = \|\mathcal{A}\mathbf{f} - \mathbf{g}\|_2^2 + \mu_2 \left\{ \sum_i \mathcal{W}(f_i) + \omega \|\mathbf{f} - \mathbf{v}\|_2^2 \right\} \quad (4.4)$$

where $\mathbf{v} \in \mathbb{R}^{N'}$, the constant $\omega = \frac{1}{2}(\frac{s_1 - s_2}{2})^2 + 1 \in \mathbb{R}^+$ is chosen so that $\mathcal{J}(\mathbf{f}|\omega, \mathbf{v})$ is convex. With the TV regularization, the final convex optimization problem is formulated

as:

$$\min_{\mathbf{f} \in \mathbb{R}^{N'}, \mathbf{u} \in \mathbb{R}^{3N'}} \mathcal{J}(\mathbf{f}|\omega, \mathbf{v}) + \mu_1 \sum_i \|\mathbf{u}_i\| \quad (4.5)$$

with the linear constraint of TV regularization $\mathcal{D}_i \mathbf{f} = \mathbf{u}_i$. $\mathbf{u}$ is a splitting variable to optimize the nonsmooth regularization term. This problem is solved by an augmented Lagrangian where $\beta \in \mathbb{R}^+$ and $\boldsymbol{\lambda} \in \mathbb{R}^{3N'}$ are Lagrangian parameter and Lagrangian multiplier respectively (Artina et al., 2013). The Lagrangian parameter controls the weight of the quadratic additional penalty. The algorithm searches for a saddle point of the augmented Lagrangian and updates the quadratic convexifying penalty. Before giving the entire algorithm, here we define the augmented Lagrangian functional

$$\begin{aligned} \mathcal{L}_A(\mathbf{f}|\omega, \mathbf{f}^{l-1}) = & \mathcal{J}(\mathbf{f}|\omega, \mathbf{f}^{l-1}) + \mu_1 \sum_i (\|\mathbf{u}_i^{(l,k-1)}\| \\ & - \langle \boldsymbol{\lambda}_i^{(l,k-1)}, \mathcal{D}_i \mathbf{f} - \mathbf{u}_i^{(l,k-1)} \rangle + \frac{\beta}{2} \|\mathcal{D}_i \mathbf{f} - \mathbf{u}_i^{(l,k-1)}\|_2^2) \end{aligned} \quad (4.6)$$

where l is the outer loop iteration number, k denotes the inner loop iteration number. A gradient descent method is applied to solve the optimization of problem:

$$\mathbf{f}^{(l,k)} = \arg \min_{\mathbf{f}} \mathcal{L}_A(\mathbf{f}|\omega, \mathbf{f}^{l-1}) \quad (4.7)$$

For $p \geq 1$:

$$\begin{aligned} \mathbf{f}^{(l,k,0)} &= \mathbf{f}^{(l,k-1)}, \quad p = 1; \\ \mathbf{f}^{(l,k,p)} &= \mathbf{f}^{(l,k,p-1)} - \alpha \{ \mathcal{A}^t(\mathcal{A} \mathbf{f}^{(l,k,p-1)} - \mathbf{g}) + \\ & \mu_1 (-\mathcal{D}^t \boldsymbol{\lambda}^{(l,k-1)} + \beta \mathcal{D}^t (\mathcal{D} \mathbf{f}^{(l,k,p-1)} - \mathbf{u}^{(l,k-1)})) + \\ & \mu_2 (\mathcal{W}'(\mathbf{f}^{(l,k,p-1)}) + 2\omega(\mathbf{f}^{(l,k,p-1)} - \mathbf{f}^{(l-1)})) \} \end{aligned} \quad (4.8)$$

where α is a small descent parameter. A saddle point is obtained with an external loop on the index l and with an internal loop on the index $1 \leq k \leq L^l$ with the following algorithm:

Algorithm 1 LCNM's algorithm

loop on l :

$$\mathbf{f}^{(l,0)} = \mathbf{f}^{l-1} := \mathbf{f}^{(l-1,L^{l-1})}$$

$$\boldsymbol{\lambda}^{(l,0)} = \boldsymbol{\lambda}^{l-1} := \boldsymbol{\lambda}^{(l-1,L^{l-1})}$$

$$\mathbf{u}^{(l,0)} = \mathbf{u}^{l-1} := \mathbf{u}^{(l-1,L^{l-1})}$$

loop on k :

while $(1 + \|\boldsymbol{\lambda}^{l-1}\|) \|\mathcal{D} \mathbf{f}^{(l,L^l)} - \mathbf{u}^{(l,L^l)}\| \leq \frac{1}{l^\gamma}$ **do**

$$\mathbf{f}^{(l,k)} = \arg \min_{\mathbf{f}} \mathcal{L}_A(\mathbf{f}|\omega, \mathbf{f}^{l-1})$$

$$\mathbf{u}^{(l,k)} = \mathcal{S}_\beta(\mathcal{D} \mathbf{f}^{(l,k)} - \frac{\boldsymbol{\lambda}^{(l,k-1)}}{\beta})$$

$$\boldsymbol{\lambda}^{(l,k)} = \boldsymbol{\lambda}^{(l,k-1)} + \beta(\mathbf{u}^{(l,k)} - \mathcal{D} \mathbf{f}^{(l,k)})$$

end while

end loop

end loop

The idea of this algorithm is to add a quadratic term to a nonconvex functional so that the optimization functional turns to be convex. The outer loop is used to update the quadratic term, more precisely it is $\mathbf{f}^{l-1}$ that is updated, and $\mathbf{f}^{l-1}$ corresponds to the variable $\mathbf{v}$ in Equation 4.4. The inner loop is used to minimize the optimization functional which is nonconvex if no quadratic term is added.

Algorithm 1 describes the algorithm of LCNM, where $\mathcal{S}_\beta$ is the soft thresholding operator, defined as $\mathcal{S}_\beta(\mathbf{x}) = \max(1 - \frac{1}{\beta|\mathbf{x}|+\epsilon}, 0)\mathbf{x}$. ϵ is a small number ensuring the stability of algorithm if $|x| \simeq 0$, γ is a constant number. L^l is the smallest number such that the following condition holds:

$$(1 + \|\boldsymbol{\lambda}^{l-1}\|)\|\mathcal{D}\mathbf{f}^{(l,L^l)} - \mathbf{u}^{(l,L^l)}\| \leq \frac{1}{l^\gamma} \quad (4.9)$$

4.3.2 Accelerated Proximal Gradient (APG)

The proximal operator and the accelerated proximal gradient have been introduced in subsection 2.4.1. Here we would like to show how to use APG to solve our nonconvex and nonsmooth optimization problem in Eq.4.5.

First of all, the objective functional $\mathcal{J}(\mathbf{f})$ is split as the sum of two functionals:

$$\begin{aligned} \mathcal{J}(\mathbf{f}) &= \mathcal{F}_1(\mathbf{f}) + \mathcal{F}_2(\mathbf{f}) \\ \mathcal{F}_1(\mathbf{f}) &= \|\mathcal{A}\mathbf{f} - \mathbf{g}\|^2 + \mu_2 \sum_i \mathcal{W}(\mathbf{f}_i) \\ \mathcal{F}_2(\mathbf{f}) &= \mu_1 \sum_i \|\mathcal{D}_i \mathbf{f}\|_2 \end{aligned} \quad (4.10)$$

where $\mathcal{F}_1$ is differentiable but nonconvex, and $\mathcal{F}_2$ is nonsmooth but convex.

Let L be the Lipschitz constant of $\nabla \mathcal{F}_1$, $\mathbf{f}^{(k)}, \mathbf{h}^{(k)}, \mathbf{z}^{(k)} \in \mathbb{R}^{N'}$, $c^{(k)}$ is a weighted average energy function. The algorithm in (Li and Lin, 2015) can be written in our case as:

This algorithm is also called non monotonous APG, since the $\mathcal{J}(\mathbf{f}^{(k+1)})$ can be greater than $\mathcal{J}(\mathbf{f}^{(k)})$, but $c^{(k+1)}$ is supposed to be smaller than $c^{(k)}$. In another word, it is $c^{(k)}$ which is minimized instead of the function $\mathcal{J}(\mathbf{f}^{(k)})$. If $\mathbf{h}^{(k)}$ is a good extrapolation point, then $\mathbf{f}^{(k+1)} = \mathbf{z}^{(k+1)}$, which accelerates the algorithm; otherwise $\mathbf{v}^{(k)}$ will be updated with a smaller step length. The proximal operator is implemented by an ADMM minimization method (Toma et al., 2014a).

4.3.3 A combination of Cahn-Hilliard model and Total Variation (CH+TV)

We finally present a third method following the work of Bertozzi *et al.* considering a fourth order inpainting algorithm for binary images based on Cahn-Hilliard equation. We first recall this method and propose an adaptation to our problem.

The Cahn-Hilliard potential in image processing tasks

The method proposed by Bertozzi, Esedoğlu and Gilette (Bertozzi et al., 2007a, Bertozzi et al., 2007b) considers the Sobolev H^{-1} gradient and convexity splitting. The H^{-1} norm

Algorithm 2 APG's algorithm

Initialization: $\mathbf{z}^{(1)} = \mathbf{f}^{(1)} = \mathbf{f}^{(0)}$, $t^{(1)} = 1$, $t^{(0)} = 0$, $\eta \in [0, 1)$, $\delta > 0$, $c^{(1)} = \mathcal{J}(\mathbf{f}^{(1)})$,
 $q^{(1)} = 1$, $\alpha_f < 1/L$, $\alpha_h < 1/L$
while $k \leq T$ **do** ▷ T: total number of iteration
 $\mathbf{h}^{(k)} = \mathbf{f}^{(k)} + \frac{t^{k-1}}{t^{(k)}}(\mathbf{z}^{(k)} - \mathbf{f}^{(k)}) + \frac{t^{k-1} - 1}{t^{(k)}}(\mathbf{f}^{(k)} - \mathbf{f}^{(k-1)})$
 $\mathbf{z}^{(k+1)} = \text{prox}_{\alpha_h \mathcal{F}_2}(\mathbf{h}^{(k)} - \alpha_h \nabla \mathcal{F}_1(\mathbf{h}^{(k)}))$
if $\mathcal{J}(\mathbf{z}^{(k+1)}) \leq c^{(k)} - \delta \|\mathbf{z}^{(k+1)} - \mathbf{h}^{(k)}\|^2$ **then**
 $\mathbf{f}^{(k+1)} = \mathbf{z}^{(k+1)}$
else
 $\mathbf{v}^{(k+1)} = \text{prox}_{\alpha_f \mathcal{F}_2}(\mathbf{f}^{(k)} - \alpha_f \nabla \mathcal{F}_1(\mathbf{f}^{(k)}))$
 $\mathbf{f}^{(k+1)} = \begin{cases} \mathbf{z}^{(k+1)}, & \text{if } \mathcal{J}(\mathbf{z}^{(k+1)}) \leq \mathcal{J}(\mathbf{v}^{(k+1)}), \\ \mathbf{v}^{(k+1)}, & \text{otherwise.} \end{cases}$
end if
 $t^{(k+1)} = \frac{\sqrt{4(t^{(k)})^2 + 1} + 1}{2}$
 $q^{(k+1)} = \eta q^{(k)} + 1$
 $c^{(k+1)} = \frac{\eta q^{(k)} c^{(k)} + \mathcal{J}(\mathbf{f}^{(k+1)})}{q^{(k+1)}}$
end while

is defined as $\|\cdot\|_{-1} = \|\nabla \Delta^{-1} \cdot\|$, where Δ^{-1} denotes the inverse of the Laplacian. They showed that the partial differential diffusion equation giving the solution of inpainting problems admits a weak solution in $H_1(\Omega)$, and their implementation is based on convexity splitting methods. The Cahn-Hilliard equation originates from diffusion driven by the chemical potential (Elliott, 1989). The concave part of the potential induces a diffusion towards the wells of the potential.

Bertozzi *et al.* (Bertozzi et al., 2007b) proposed a modified Cahn-Hilliard equation to solve some inpainting problems. Given an observation image $\mathbf{g}$, the evolution equation for the restored image $\mathbf{f}$ is the following:

$$\mathbf{f}_t = -\Delta(\xi \Delta \mathbf{f} - \frac{1}{\xi} \mathcal{W}'(\mathbf{f})) + \mathbf{m} \odot (\mathbf{g} - \mathbf{f}) \quad (4.11)$$

$\mathbf{m}$ serves as a mask preserving the region to be inpainted, and $\odot$ is element-wise multiplication. The parameter ξ controls the diffusion process: depending on its value, structures may be connected roughly or refined delicately. Besides the inpainting problem, (Bertozzi et al., 2007b) also tries to solve some super-resolution problems with the same modified Cahn-Hilliard model.

Algorithm of CH+TV

In this section, we combine the inpainting technique in (Bertozzi et al., 2007b) with the Cahn-Hilliard model and the TV regularization. The objective functional considered is:

$$\mathcal{J}(\mathbf{f}, \mathbf{u}_i, \mathbf{v}_i) = \frac{1}{2} \|\mathcal{A}\mathbf{f} - \mathbf{g}\|^2 + \sum_i \{\mu_1 \|\mathbf{u}_i\| + \mu_2 (\mathcal{W}(\mathbf{v}_i)/\xi + \xi \|\nabla \mathbf{v}_i\|^2)\} \quad (4.12)$$

with linear constraints $\forall i, \mathbf{u}_i = \mathcal{D}_i \mathbf{f}, \mathbf{v} = \mathbf{f}$. ξ is a hyper-parameter. The augmented Lagrangian including these constraints is given by:

$$\begin{aligned} \mathcal{L}_A(\mathbf{f}, \mathbf{u}_i, \mathbf{v}, \boldsymbol{\lambda}_i, \boldsymbol{\lambda}_D) = & \frac{1}{2} \|\mathcal{A}\mathbf{f} - \mathbf{g}\|_2^2 \\ & + \mu_1 \sum_i \{ \|\mathbf{u}_i\| + \frac{\beta}{2} \|\mathbf{u}_i - \mathcal{D}_i \mathbf{f}\|^2 - \boldsymbol{\lambda}_i^t (\mathbf{u}_i - \mathcal{D}_i \mathbf{f}) \} \\ & + \mu_2 \{ \sum_i (\frac{1}{\xi} \mathcal{W}(\mathbf{v}_i) + \xi \|\nabla \mathbf{v}_i\|^2) + \frac{\beta}{2} \|\mathbf{v} - \mathbf{f}\|_2^2 - \boldsymbol{\lambda}_D^t (\mathbf{v} - \mathbf{f}) \} \end{aligned}$$

$\boldsymbol{\lambda}_i$ and $\boldsymbol{\lambda}_D$ are Lagrangian multipliers for TV and double-well functions respectively. A saddle point of Lagrangian function is estimated with successive updates.

1. Update of $\mathbf{f}^{(k+1)}$:

$$\begin{aligned} & (\mathcal{A}^t \mathcal{A} + \mu_1 \sum_i \beta \mathcal{D}_i^t \mathcal{D}_i + \mu_2 \beta \mathbf{I}) \mathbf{f}^{(k+1)} \\ & = \mathcal{A}^t \mathbf{g} + \mu_1 \sum_i (\mathcal{D}_i^t (\beta \mathbf{u}_i^{(k)} - \boldsymbol{\lambda}_i^{(k)})) + \mu_2 (\beta \mathbf{v}^{(k)} - \boldsymbol{\lambda}_D^{(k)}) \end{aligned} \quad (4.13)$$

The solution of this linear equation is obtained with a conjugate gradient method.

2. Update of $\mathbf{u}_i^{(k+1)}$:

$$\mathbf{u}_i^{(k+1)} = \mathcal{S}_\beta(\mathcal{D}_i \mathbf{f}^{(k+1)} + \boldsymbol{\lambda}_i^{(k)} / \beta) \quad (4.14)$$

3. Update of $\mathbf{v}^{(k+1)}$, here we use the approach of Esedoğlu *et al.* in (Bertozzi et al., 2007b), which is based on the H^{-1} norm.

$$\begin{aligned} & \mathbf{v}^{(k+1)} + \Delta t (\xi \Delta^2 \mathbf{v}^{(k+1)} - C_1 \Delta \mathbf{v}^{(k+1)} + C_2 \mathbf{v}^{(k+1)}) \\ & = \mathbf{v}^{(k)} + \Delta t (\Delta (\frac{1}{\xi} \mathcal{W}'(\mathbf{v}^{(k)})) + \beta (\mathbf{f}^{(k+1)} - \mathbf{v}^{(k)}) + \\ & \quad \boldsymbol{\lambda}_D - C_1 \Delta \mathbf{v}^{(k)} + C_2 \mathbf{v}^{(k)}) \end{aligned} \quad (4.15)$$

where Δt is time step length. Following the work in (Bertozzi et al., 2007b, Eyre, 1998), the solution can be obtained by Fourier transforms. The constants C_1 , C_2 and ξ should be well chosen to ensure the stability of this algorithm.

4. Update of the Lagrange multipliers $\boldsymbol{\lambda}_i^{(k+1)}$ and $\boldsymbol{\lambda}_D^{(k+1)}$:

$$\begin{aligned} \boldsymbol{\lambda}_i^{(k+1)} &= \boldsymbol{\lambda}_i^{(k)} - \beta (\mathbf{u}_i^{(k+1)} - \mathcal{D}_i \mathbf{f}^{(k+1)}) \\ \boldsymbol{\lambda}_D^{(k+1)} &= \boldsymbol{\lambda}_D^{(k)} - \beta (\mathbf{v}^{(k+1)} - \mathbf{f}^{(k+1)}) \end{aligned} \quad (4.16)$$

To obtain stable numerical schemes, some splitting methods for the gradient flow of the double well potential must be used. (Bertozzi et al., 2007b) refers to (Eyre, 1998) to split the Cahn-Hilliard diffusion into contractive functions and expansive functions. Since the

modified Cahn-Hilliard equation consists of double well functional (E_1) and data fitting term (E_2), its energy can be written as:

$$\begin{aligned} E &= E_1 + E_2 \\ E_1(\mathbf{v}) &= \int_{\Omega} \frac{\xi}{2} |\nabla \mathbf{v}|^2 + \frac{1}{\xi} \mathcal{W}(\mathbf{v}) dx \\ E_2(\mathbf{v}) &= \lambda_0 \int_{\Omega} (\mathbf{f} - \mathbf{v})^2 dx \end{aligned} \tag{4.17}$$

where E_1 is minimized by H^{-1} gradient flow, and E_2 with a L^2 gradient flow. The splitting of E_1 and E_2 is detailed in the following:

$$\begin{aligned} E_1 &= E_{11} - E_{12} \\ E_{11} &= \int_{\Omega} \frac{\xi}{2} |\nabla \mathbf{v}|^2 + \frac{C_1}{2} |\mathbf{v}|^2 dx \\ E_{12} &= \int_{\Omega} -\frac{1}{\xi} \mathcal{W}(\mathbf{v}) + \frac{C_1}{2} |\mathbf{v}|^2 dx \end{aligned} \tag{4.18}$$

Similarly for E_2 , we write:

$$\begin{aligned} E_2 &= E_{21} - E_{22} \\ E_{21} &= \int_{\Omega} \frac{C_2}{2} |\mathbf{v}|^2 dx \\ E_{22} &= \int_{\Omega} -\frac{\beta}{2} (\mathbf{f} - \mathbf{v})^2 + \frac{C_2}{2} |\mathbf{v}|^2 dx \end{aligned} \tag{4.19}$$

The constant C_1 and C_2 must be chosen large enough so that the energies E_{11} , E_{12} , E_{21} and E_{22} are convex to ensure the convergence of the algorithm.

4.4 Application to bone CT images

4.4.1 Methodology

The tests were performed on two low resolution experimental images of trabecular bone samples obtained from HR-pQCT (Boutroy et al., 2005) with a voxel size of $82\mu\text{m}$. And they will be denoted as (A) and (B) in the following. Meanwhile, two high resolution images were also obtained on the same bone samples from μCT with a voxel size of $41\mu\text{m}$ after a registration. The registration of the experimental low and high resolution images was performed with an open source software Elastix (Shamonin et al., 2014, Klein et al., 2010) with B-spline interpolation and stochastic gradient descent method. The blurring operator we applied in this work has been determined in a previous work (Li et al., 2017b).

All computations were performed on the 3D volumes. We processed cropped ROIs (region of interest) of size $100 \times 100 \times 100$ and $200 \times 200 \times 200$ in the low and high resolution images respectively. A first restoration approach based on Total Variation (TV) regularization was used as a starting point for the three minimization methods. This method was developed in a previous work and is detailed in (Toma et al., 2014a, Toma

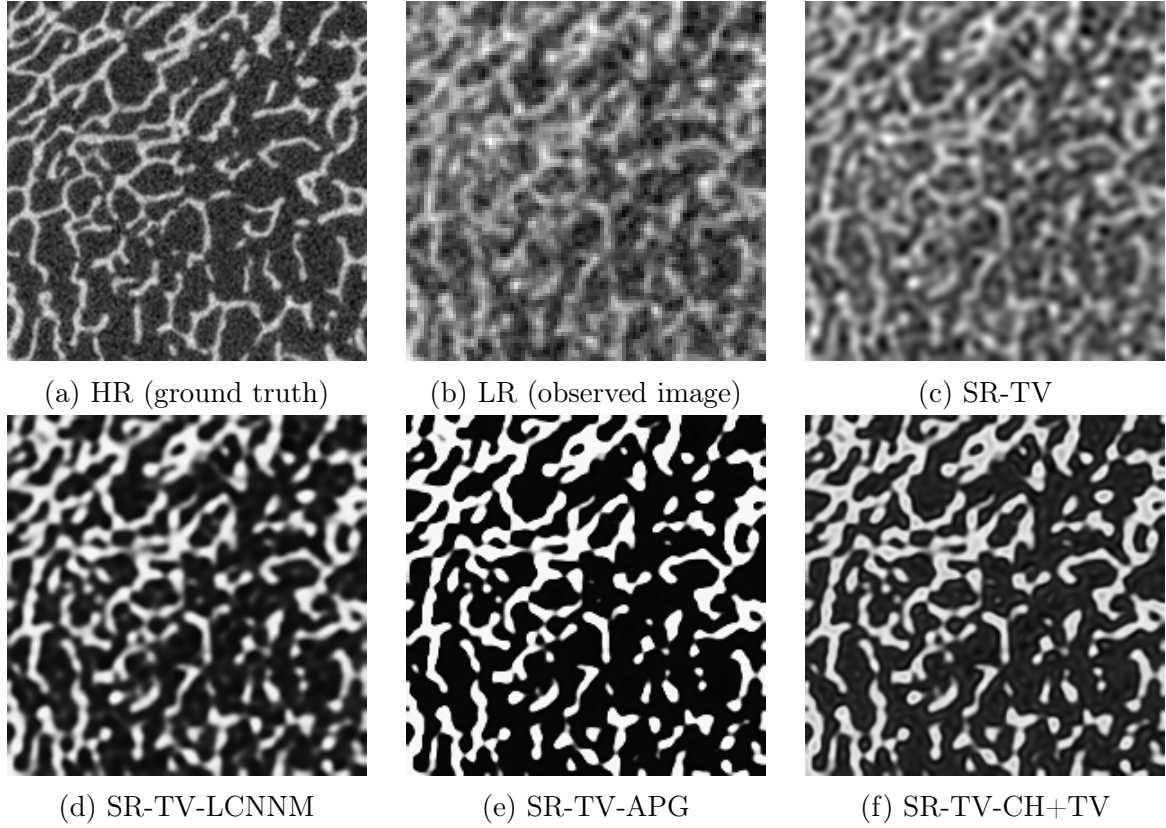


Figure 4.3: 2D slices from the first crop (A): HR and LR denote high resolution image and low resolution image respectively. SR-TV is the super resolution image restored by TV. SR-TV-LCNNM, SR-TV-APG and SR-TV-CH+TV represent images reconstructed by LCNM, APG and CH+TV initialized with SR-TV.

et al., 2015). The regularization functionals are non convex and the results obtained with non regularized starting conditions are very poor.

In order to evaluate the quality of the images restored with the different algorithms, we used several criteria. The assessment was performed on three images: $\mathbf{f}$, $\mathbf{f}_b$ and $\mathbf{f}_p$, corresponding to the gray level image, the binarized image (segmented with the Otsu(Otsu, 1975) method) and the image obtained by separately averaging the gray level image $\mathbf{f}$ over the different segmented regions on $\mathbf{f}_b$. $\mathbf{f}_p$ is obtained as follows:

1. Given a resolved image $\mathbf{f}$, its binary image is $\mathbf{f}_b$. $\mathbf{f}_b$ only has 0 and 1.
2. If we regard $\mathbf{f}_b$ as a mask of $\mathbf{f}$, we could calculate the average value $\mathbf{f}_0$ and $\mathbf{f}_1$ of $\mathbf{f}$ over zeros regions and ones regions of $\mathbf{f}_b$.
3. The voxels of value 0 in $\mathbf{f}_b$ is assigned to $\mathbf{f}_0$, the voxel of value 1 in $\mathbf{f}_b$ is set to $\mathbf{f}_1$.

We used the classical PSNR measure, the DICE, the BV/TV and the Conn.D as criteria to qualify the reconstructions.

The DICE, BV/TV and Conn.D have been calculated on the segmented images at each iteration. The projected data term was estimated by $\|\mathbf{A}\mathbf{f}_p - \mathbf{g}\|_2^2$. This term is a good

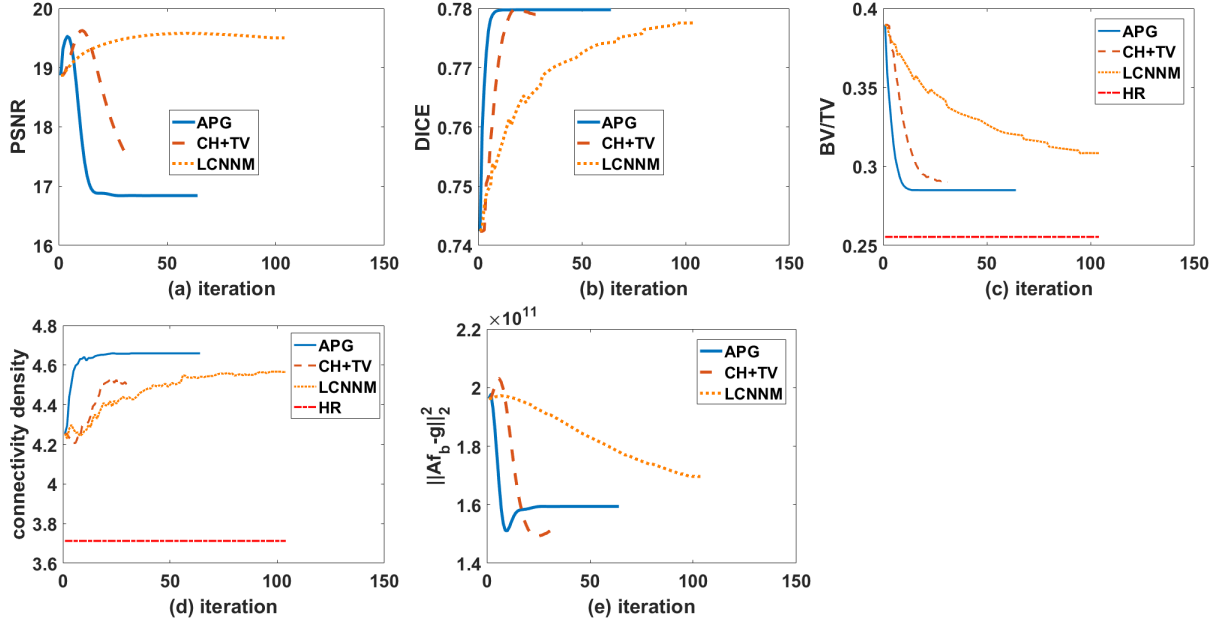


Figure 4.4: For the first crop (A): (a)-(e) illustrate the evolutions of PSNR, DICE, BV/TV, Conn.D and $\|\mathcal{A}f_p - g\|_2^2$ as a function of the iteration number. For BV/TV and Conn.D, the high resolution value is also displayed as a dashed straight line on the graphs.

indicator of the quality of the restoration and it is used to obtain the best approximation of the high resolution image, as explained in section 4.4.2.

The regularization parameter values were chosen after an extensive sweeping to obtain the best decrease of the projected data term. For LCNM, the gradient descent parameter is $\alpha = 0.01$, the Lagrangian parameter is $\beta = 500$, $\epsilon = 10^{-6}$ so that the shrink operator $\mathcal{S}_\beta$ is stable. We set $\mu_1 = 0.1$ and $\mu_2 = 10^{-9}$ for the weights of the TV regularization and the double-well potential. In APG algorithm, the weights for the regularization terms are $\mu_1 = 100$ and $\mu_2 = 2 \times 10^{-9}$, the Lagrangian parameter is $\beta = 500$, the scale factors in the proximal operator are set to be $\alpha_h = 1$ and $\alpha_f = 0.5$, $\eta = 0.8$, $\delta = 10$. As for CH+TV, the weights of the TV regularization and the Cahn-Hillard model are $\mu_1 = 15$ and $\mu_2 = 0.5$ respectively, the Lagrangian parameter is $\beta = 500$, $\xi = 1$, $C_1 = 8 \times 10^8$, $C_2 = 500$.

4.4.2 Results

Figs.4.3 and 4.5 show 2D displays of the high and low resolution images, the initial image obtained with TV, the reconstructed images obtained by the three approaches discussed above. For the LCNM, we display the last iterate, the APG and the CH+TV images were obtained with an early stopping iterate. As expected, the images reconstructed by the three methods are finally quasi binary. Figs.4.4 and 4.6 present the evolution of the image quality metrics for the different algorithms. Figs.4.4(a) and 4.6(a) present the evolution of PSNR as a function of the iteration number for two crops (A) and (B). CH+TV and APG are similar in that their PSNR increase and then decline, while LCNM reaches a plateau after a steady increase. As illustrated in Fig.4.4(b) and 4.6(b), these

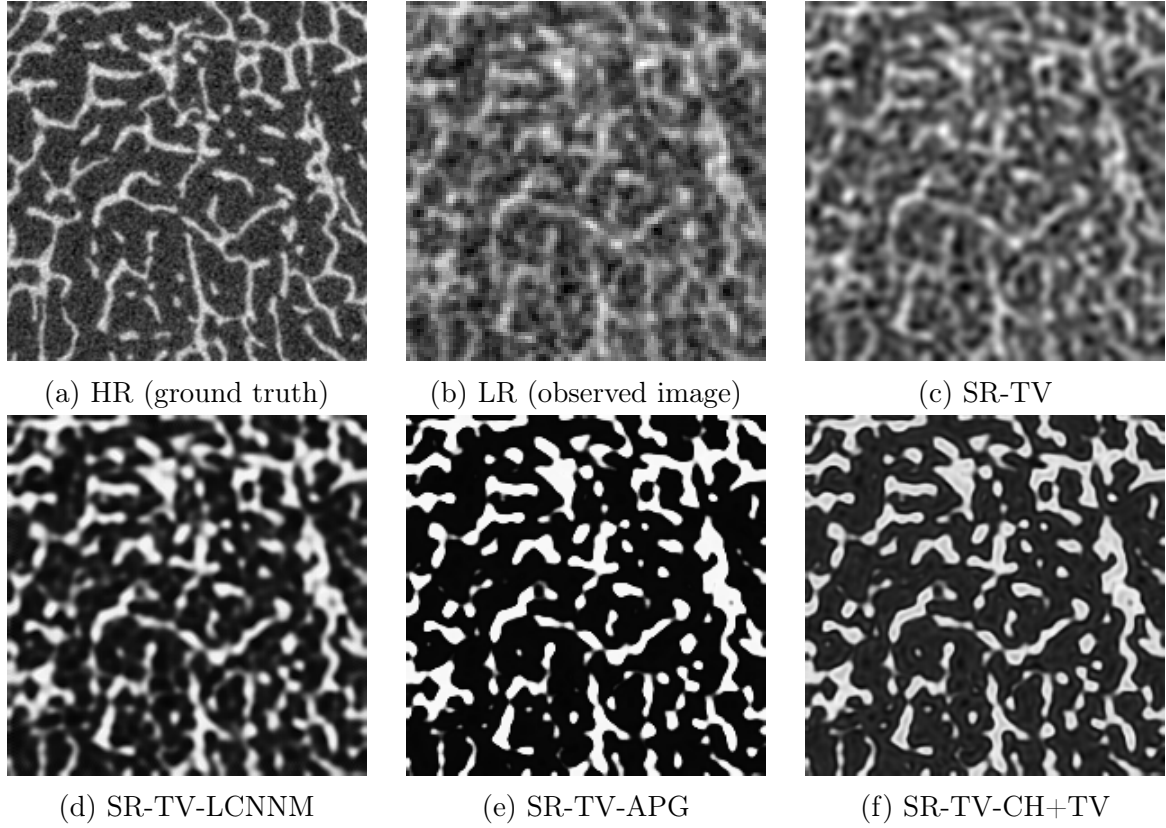


Figure 4.5: 2D slices from the second crop (B): HR and LR denote high resolution image and low resolution image respectively. SR-TV is the super resolution image restored by TV. SR-TV-LCNNM, SR-TV-APG and SR-TV-CH+TV represent images reconstructed by LCNM, APG and CH+TV initialized with SR-TV.

three methods improve the DICE of the image. As for Fig.4.4(c) and 4.6(c), BV/TV is improved by these three methods. In addition, both APG and CH+TV drop sharply while LCNM declines gradually. However, in Fig.4.4(d), the Conn.D obtained by these three approaches rises instead of approximating the Conn.D of the high resolution image. A steep increase of Conn.D can be interpreted as the deterioration of the bone structure. A small improvement in the Conn.D is obtained for the second crop (B) with LCNM scheme. No improvement is observed for the other algorithms.

As expected, the image estimates are going towards different local minima. And on these plots, some of these algorithms have not really converged. Yet, we can see that an early truncation of the algorithm may be useful to obtain a better image. We have not stopped the algorithm at the maximum PSNR since PSNR is computed with the ground truth, which is unknown in practice. On contrary, a projected data term is determined without the knowledge of ground truth. Fig.4.4(e) and Fig.4.6(e) present the evolution with the iterations of the projected data term $\|\mathcal{A}\mathbf{f}_p - \mathbf{g}\|_2^2$. For CH+TV, the iterations are stopped at the minimum of $\|\mathcal{A}\mathbf{f}_p - \mathbf{g}\|_2^2$. The image obtained with this stopping rule has similar or even better image quality criteria than the one with the best PSNR, see Fig.4.4 and Fig.4.6. This iteration stopping strategy was also applied to APG. For these two methods, the stopping rule leads to a solution near the best values of the structural

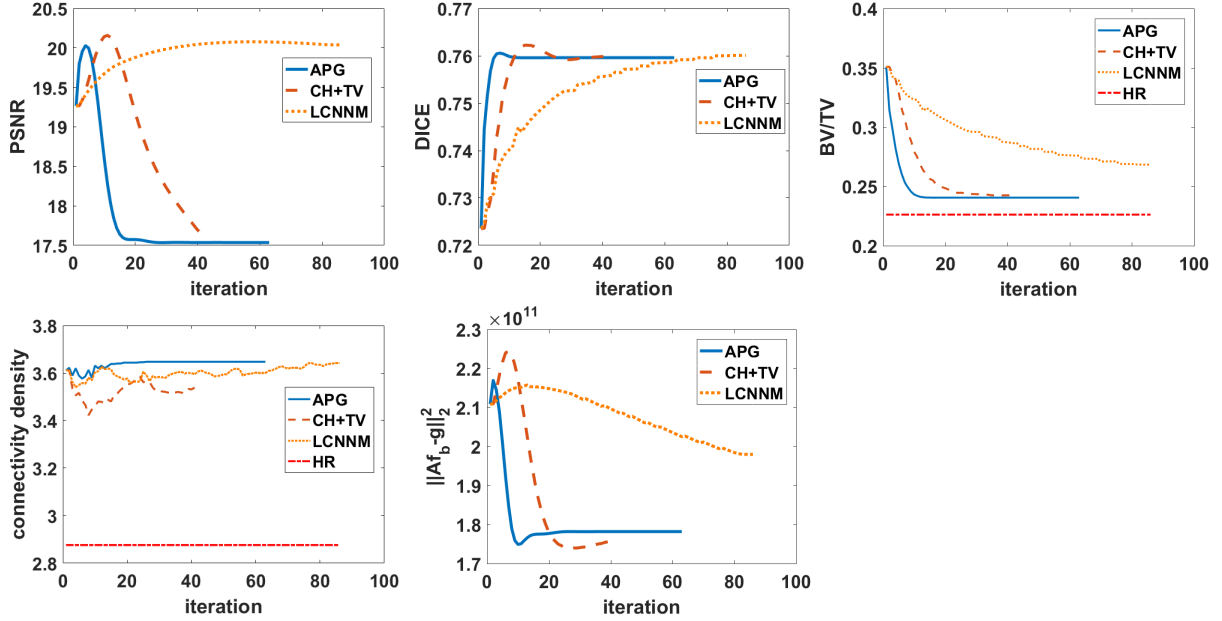


Figure 4.6: For the second crop (B): (a)-(e) illustrate the evolutions of PSNR, DICE, BV/TV, Conn.D and $\|\mathcal{A}f_p - g\|_2^2$ as a function of the iteration number. For BV/TV and Conn.D, the high resolution value is also displayed as a dashed straight line on the graphs.

and topological parameters. However, this does not change the effect of the deterioration of Conn.D by APG.

4.5 Discussion and Conclusions

In this work, we combined TV and a double-well potential to solve a joint super resolution/segmentation problem. All the three algorithms described above can be used to minimize nonconvex (double-welled potential) and nonsmooth (TV) functional. TV regularization term can be minimized by ADMM which could also be regarded as the implementation of proximal operator for the dual problem. The LCNNM minimizes nonconvex functional by introducing a quadratic term so that nonconvex term $\mathcal{W}$ turns to be convex. As for the APG, extrapolation points are used to correct the descent direction and accelerate the convergence of the algorithm.

The aim of this work is to facilitate the segmentation and ameliorate the evaluation of the bone structural parameters. We consider that the DICE, BV/TV and Conn.D which are calculated on the segmented images are more relevant to evaluate the efficiency of the method than PSNR. They show how the restoration method improves the subsequent segmentation which is performed with the Otsu method.

From a convergence perspective, the APG method wins. LCNNM is less efficient since it converges very slowly and leads to a smaller improvement in the BV/TV. LCNNM would be the best one with a large number of iterations. Yet, with an early truncation of the algorithm based on $\|\mathcal{A}f_p - g\|_2^2$, the conclusion is reversed. APG improves the DICE

and the BV/TV but degrades the connectivity. The Conn.D is a topological parameter that is crucial for the evaluation bone strength. We consider that CH+TV is the best one since it degrades less the connectivity with similar improvements in the DICE and of the BV/TV than APG.

The images reconstructed by the three approaches have better contrast due to double-well potential and a smooth structure in response to TV. Due to the non convexity of the regularization functional, similar structures with different morphology and topology parameters are obtained with the different algorithms. We have shown that an early stopping of the algorithms may be useful to obtain the best structural and topological parameters. The CH+TV is considered as the best method among the three algorithms investigated. For this algorithm, criteria such as DICE and BV/TV are improved. Yet, the proposed methods do not preserve the connectivity of the structure.

Mathematically speaking, the TV regularization or Laplacian operator takes account of very small neighbor regions. The double-welled potential adds prior information on the intensity of each voxels. The exploration of image structural information is still limited to small regions.

Fig.4.3 and Fig.4.5 show that the reconstructed super resolution images lost many details, compared with the ground truth. In the following, we intend to add more prior about the trabeculea structure for the purpose of recovering more details. Dictionary learning approach is a good candidate because it allows to extract images characteristic within a larger region than TV or CH+TV. Dictionary learning approaches have been introduced in 2.2.4. We're going to introduce the application of the dictionary learning in our problem in the next chapter.

Chapter 5

Investigation of semi-coupled dictionary learning in 3D super resolution HR-pQCT imaging

5.1 Introduction

Dictionary learning is another approach to solve super resolution problems. A joint dictionary learning method has been studied by Yang *et al.* (Yang *et al.*, 2010a). The same authors have proposed a coupled-dictionary learning method (Yang *et al.*, 2012). The principle of coupled-dictionary learning is to couple two databases to update dictionary atoms. Wang *et al.* have also investigated a semi-coupled dictionary learning (Wang *et al.*, 2012)(SCDL). In this case, mapping matrices for the sparse coefficients of the high and low resolution patches are learned. He *et al.* used a Bayesian method and a Beta process for coupled dictionary learning(He *et al.*, 2013). Some authors have considered locality constraints(Wang *et al.*, 2012, He *et al.*, 2013, Wang *et al.*, 2018).

Dictionary learning has been previously used in medical image processing, mostly for denoising(Li *et al.*, 2012), compressed sensing(Huang *et al.*, 2014) or low dose reconstruction(Xu *et al.*, 2012, Chen *et al.*, 2013, Soltani *et al.*, 2017). Joint dictionary learning has been investigated for Optical Coherence Tomography(OCT) (Asif *et al.*, 2017)(Asif *et al.*, 2016) and microscopy image registration (Cao *et al.*, 2014). Coupled dictionary learning has, for instance, been used to reduce streak artifacts in CT reconstruction (Karimi and Ward, 2016) or to correct MRI image deformation (Cao *et al.*, 2015).

In this work, we investigate an approach based on SCDL (Wang *et al.*, 2012) to increase the resolution of 3D HR-pQCT images based on the knowledge of sets of low and high resolution images. To the best of our knowledge, this is the first work to investigate SCDL for experimental CT images with a complex architecture. Moreover, the application of this approach in 3D and the comparison with other methods based on TV or TV coupled with a double well potential is original. The processing of 3D images also adds complexity in the process with respect to the 2D case investigated in the original work (Wang *et al.*, 2012) because of the increase of memory requirements and computing time. Actually, the complexity of SCDL tends to exponentially increase with respect to the number of dimensions of patches, which hinders the investigation of 3D SCDL. Nevertheless, the characterization of connectivity and topology properties of bone requires 3D images. In

order to save computing time and to handle bone anisotropy, we propose a 2.5D strategy to learn high and low resolution dictionaries on coronal/sagittal/axial directions separately.

This chapter is organized as follows. In section 5.2, we introduce the SCDL approach, then explain the proposed strategies to process our 3D images. In section 5.3, the results on experimental images are presented and compared with previously developed methods based on TV regularization. Some discussions about the comparative efficiency of the methods are concluding the chapter.

5.2 Semi-coupled dictionary learning approach

The semi-coupled dictionary learning method(Wang et al., 2012) learns two dictionaries in a semi-coupled way, one dictionary for high resolution images, denoted as Φ^f , the other for low resolution images, denoted as Φ^g . In the following, we will denote the size of the patches and the number of patches as P and Q respectively: for $1 \leq i \leq Q$, $f_i \in \mathbb{R}^{P \times 1}$ is one patch in the high resolution dataset, $g_i \in \mathbb{R}^{P \times 1}$ is a low resolution patch after interpolation. The number of atoms in Φ^f and Φ^g are K^f and K^g , hence $\Phi^f \in \mathbb{R}^{P \times K^f}$, $\Phi^g \in \mathbb{R}^{P \times K^g}$, $\alpha_i^f \in \mathbb{R}^{K^f \times 1}$ and $\alpha_i^g \in \mathbb{R}^{K^g \times 1}$ are the sparse representations of patch f_i and g_i , therefore $\alpha^f \in \mathbb{R}^{K^f \times Q}$, $\alpha^g \in \mathbb{R}^{K^g \times Q}$. In the following, we describe the two main steps of SCDL: dictionary learning and super resolution stages.

5.2.1 Dictionary learning stage

Even though g (after interpolation) and f could be sparsely represented by α^g and α^f with Φ^g and Φ^f respectively, mapping operators are needed to connect α^g and α^f . In (Wang et al., 2012), two mapping matrices W^f and W^g are introduced so that $\alpha^g = W^f \alpha^f$ and $\alpha^f = W^g \alpha^g$. The matrices W^f and W^g must be approximately inverse matrices, implying that the number of atoms in both dictionaries must be equal, i.e. $K^f = K^g = K$. We assume that the additional details in the high resolution images can be captured in the high resolution atoms and not by an increase of the number of atoms. The objective functional for the dictionary learning is given by:

$$\begin{aligned} \min_{\Phi^g, \Phi^f, W^g, W^f, \alpha^f, \alpha^g} & \|g - \Phi^g \alpha^g\|_F^2 + \|f - \Phi^f \alpha^f\|_F^2 + \\ & \gamma \|\alpha^f - W^g \alpha^g\|_F^2 + \gamma \|\alpha^g - W^f \alpha^f\|_F^2 + \\ & \lambda^W \|\mathbf{W}^g\|_F^2 + \lambda^W \|\mathbf{W}^f\|_F^2 + \\ & \lambda^f \|\alpha^f\|_1 + \lambda^g \|\alpha^g\|_1 \\ \text{s.t.} & \|d_i^g\|_2 \leq 1, \|d_i^f\|_2 \leq 1, \forall i = 1, 2, \dots, K \end{aligned} \quad (5.1)$$

where $\|\cdot\|_F$ is the Frobenius norm, d_i^g and d_i^f are the columns in matrices Φ^g and Φ^f , the index of the column is i .

1) The coefficients α^f and α^g are determined by:

$$\begin{aligned} \arg \min_{\alpha^f, \alpha^g} & \|g - \Phi^g \alpha^g\|_F^2 + \|f - \Phi^f \alpha^f\|_F^2 + \\ & \gamma \|\alpha^f - W^g \alpha^g\|_F^2 + \gamma \|\alpha^g - W^f \alpha^f\|_F^2 + \\ & \lambda^f \|\alpha^f\|_1 + \lambda^g \|\alpha^g\|_1 \end{aligned} \quad (5.2)$$

This functional is minimized using the SPAM software implementing the Lasso algorithm (Mairal et al., 2010) (Mairal et al., 2009).

2) During the training process, atoms in Φ^f and Φ^g are updated column by column (Yang et al., 2010b), which is better than a parallel updating. Let's denote η_i the i_{th} row vector of α^f , if we want to update one atom d_i^f , Equ.5.1 can be reduced to:

$$\hat{d}_i^f = \arg \min_{d_i^f} \|\mathbf{f} - \Phi^f \alpha^f\|_F^2 \quad \text{s.t.} \|\mathbf{d}_i^f\|_2 \leq 1 \quad (5.3)$$

Since $\mathbf{d}_i^f \eta_i$ represents the contribution of the atom $\mathbf{d}_i^f$ over the whole image sets $\mathbf{f}$, we use $\mathbf{Y}$ to denote the complement of $\mathbf{d}_i^f \eta_i$:

$$\mathbf{Y} = \mathbf{f} - \sum_{j \neq i} \mathbf{d}_j^f \eta_j \quad (5.4)$$

Then the problem in Equ.5.3 can be rewritten as

$$\hat{d}_i^f = \arg \min_{d_i^f} \frac{\mu}{2} \|\mathbf{Y} - \mathbf{d}_i^f \eta_i\|_F^2 - \frac{\gamma}{2} (\text{tr}(\mathbf{d}_i^f (\mathbf{d}_i^f)^T) - 1) \quad (5.5)$$

where the second term with trace operation corresponds to the norm constraint of $\mathbf{d}_i^f$. We obtain thus:

$$\begin{aligned} 0 &= \mu(\mathbf{Y} - \mathbf{d}_i^f \eta_i)(-\eta_i^T) - \gamma \mathbf{d}_i^f \\ \hat{d}_i^f &= \frac{\mu \mathbf{Y} \eta_i^T}{\mu \eta_i \eta_i^T - \gamma} \end{aligned} \quad (5.6)$$

The dictionary atom $\hat{d}_i^f$ is then replaced by its normalized value. The atoms in Φ^g are updated in the same way.

3) In the semi coupled dictionary learning, $\mathbf{W}^f$ and $\mathbf{W}^g$ can be derived by:

$$\begin{aligned} \mathbf{W}^f &= \arg \min \|\alpha^g - \mathbf{W}^f \alpha^f\|_F^2 + \lambda_W \|\mathbf{W}^f\|_F^2 \\ \mathbf{W}^g &= \arg \min \|\alpha^f - \mathbf{W}^g \alpha^g\|_F^2 + \lambda_W \|\mathbf{W}^g\|_F^2 \end{aligned} \quad (5.7)$$

λ_W is a regularization parameter. In Wang et al's work (Wang et al., 2012), $\mathbf{W}^f$ and $\mathbf{W}^g$ are updated iteratively for $k \geq 0$ by

$$\begin{aligned} \mathbf{W}_{k+1}^f &= (1 - \rho) \mathbf{W}_k^f + \rho \frac{\alpha^g \alpha^{fT}}{\alpha^f \alpha^{fT} + \lambda_W \mathbf{I}} \\ \mathbf{W}_{k+1}^g &= (1 - \rho) \mathbf{W}_k^g + \rho \frac{\alpha^f \alpha^{gT}}{\alpha^g \alpha^{gT} + \lambda_W \mathbf{I}} \end{aligned} \quad (5.8)$$

where $\rho = 0.05$ and $\mathbf{I}$ is an identity matrix.

5.2.2 Super resolution stage

Once Φ^g , Φ^f , $\mathbf{W}^g$ and $\mathbf{W}^f$ are learned, the super-resolution problem is solved by:

$$\begin{aligned} \min_{\alpha_i^f, \alpha_i^g, \mathbf{f}_i} & \|\mathbf{g}_i - \Phi^g \alpha_i^g\|_F^2 + \|\mathbf{f}_i - \Phi^f \alpha_i^f\|_F^2 + \\ & \gamma' \|\alpha_i^f - \mathbf{W}^g \alpha_i^g\|_F^2 + \gamma' \|\alpha_i^g - \mathbf{W}^f \alpha_i^f\|_F^2 + \\ & \lambda^g \|\alpha_i^g\|_1 + \lambda^f \|\alpha_i^f\|_1 \end{aligned} \quad (5.9)$$

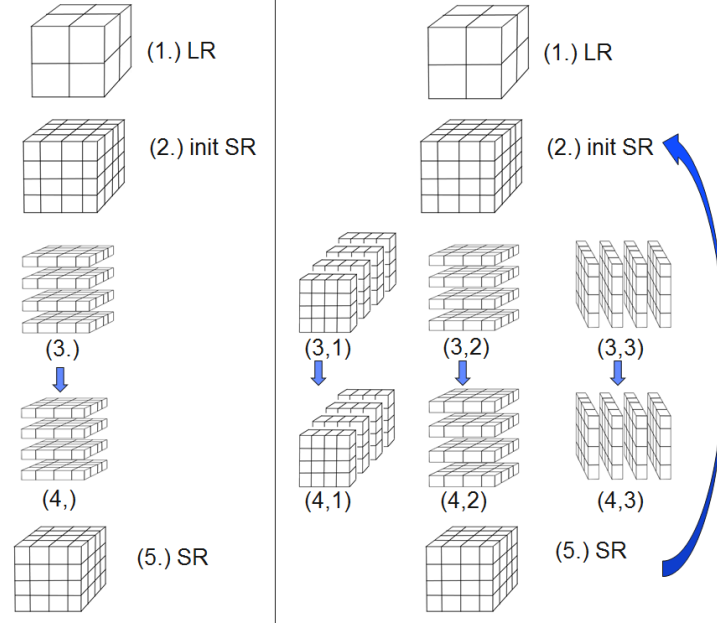


Figure 5.1: The part on the left is the flowchart of 2D SCDL for 3D volume reconstruction; the part on the right is the flowchart of 2.5D SCDL for 3D volume reconstruction. (1.)LR is the low resolution image; (2.)init SR is the initial super resolution in 3D; (3.)(3.1)(3.2)(3.3) are 2D slices cut from 3D volume in different directions (Note that (3.) and (3.2) volumes are cut horizontally); (4.)(4.1)(4.2)(4.3) are super resolution images restored with SCDL in 2D dimension; (5.) on the left is a volume stacked with slices obtained previous (4.); (5.) on the right takes the average of (4.1)(4.2)(4.3).

where $\mathbf{g}_i$ and $\mathbf{f}_i$ are i_{th} vectorized patch of $\mathbf{g}$ and $\mathbf{f}$ as explained above. The variables α_i^f , α_i^g are updated alternatively in Eq.5.9. Finally, the reconstructed patch is determined with $\hat{\mathbf{f}}_i = \Phi^f \alpha_i^f$ and the 2D super resolution image can be estimated by rearranging the resolved patches. It's noteworthy that the overlap region of patches is determined by taking the average.

5.2.3 Implementation details

The implementation of this algorithm is adapted from the codes provided by Wang *et al.* (Wang et al., 2012) via the link <http://www.comp.polyu.edu.hk/~cslzhang/SCDL.htm>. For trabecular bone imaging, we may not benefit from priors like spatial prior or nonlocal self-similarity, which is different from original SCDL application to photo sketch. As for the choice of parameters, we have swept the parameters centered with the one proposed in (Wang et al., 2012), the ratio of two successive tested parameters is 5 or 10. Then we choose the best ones by comparing their super resolution results. Furthermore, in the HR-pQCT experimental images the structural information is anisotropic with a complex noise distribution. It is not possible to suppress noise-like atoms in the dictionaries by simply reducing the number of atoms. Therefore, it's necessary to adapt the algorithm to our super resolution problem.

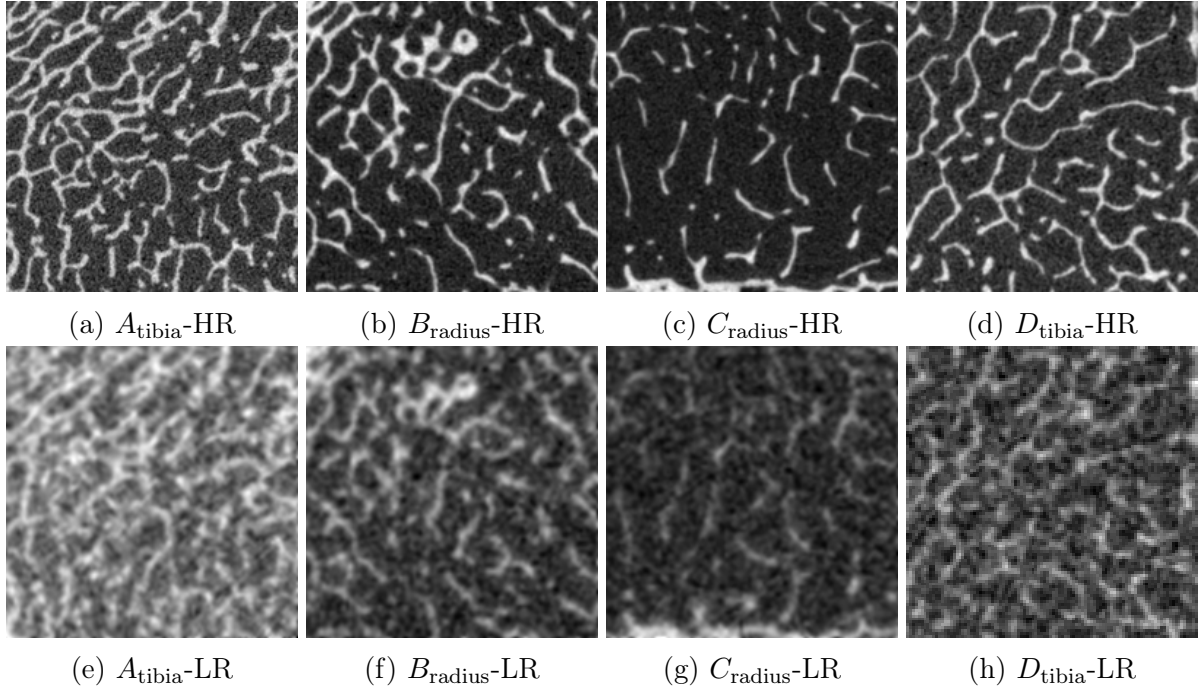


Figure 5.2: 2D slices cut from 3D samples.

5.2.4 SCDL scheme for 3D trabecular bone images

The complexity of the training stage for SCDL scales as $\mathcal{O}(p^{3c})$, where p is the edge length of patches, c is the dimensionality of patches. Thus training 3D SCDL and refining parameters turn to be a big challenge. Meanwhile, the application of SCDL in 2D slice by slice doesn't take account of the anisotropy of the bone micro-architecture. For instance, a dictionary learned from axially cut patches only contains information in axial direction.

Our strategy considers 2D dictionary learning but in 3 orthogonal directions to reconstruct 3D images, which is named as 2.5D (Luo et al., 2016) SCDL. Under this framework, the patches in the dataset are kept in 2D. This method is expected to capture the anisotropic features of the samples and to be much less time consuming than the application of the same approach in 3D. The flowchart on the right in Fig.5.1 describes the super resolution process of 2.5D SCDL. The interpolated 3D volume is cut into slices along three different directions (step 1-3), the slices will be restored with different dictionaries (step 4.1-4.3). Step 5 obtains the intermediate super resolution volume by averaging the stacks in step 4.1-4.3. Then the algorithm goes back to step 2 with the current image as starting point. In our approach, the iterative 2D SCDL is applied in each direction in order to extract the corresponding information. The iterative 2.5D SCDL improves the super resolution image by combining the information of the three directions.

5.3 Results

Imaging data were collected for a custom-built bone structure phantom (Burghardt et al., 2013). The phantom was comprised of cadaveric distal radius and tibia sections (1-cm thick) embedded in an 8-cm diameter cylinder comprised of polymethylmethacrylate

(PMMA) and polyethylene resin. The phantom was scanned using typical *in vivo* protocol settings by HR-pQCT (XtremeCT, Scanco Medical AG), which generates reconstructions with $82\mu\text{m}$ isotropic voxels. The reference scan data was collected on laboratory μCT system (Scanco Medical μCT 100), which provided high resolution reconstructions, with $24\mu\text{m}$ isotropic voxels, after registration with B-spline interpolator (Elastix (Klein et al., 2010, Shamounin et al., 2014)), its voxel size is equal to $41\mu\text{m}$.

The voxel size of the low and high resolution images are $82\mu\text{m}$ and $41\mu\text{m}$ respectively. The size of HR-pQCT images is $420 \times 510 \times 110$, and that of μCT images is $840 \times 1020 \times 220$. Seven bone samples were used in this work, three tibia and five radius. These samples are labeled from A to H. Every label has a subscript indicating its sample type. And each sample contains a couple of images. For instance, the sample A_{tibia} is a tibia bone sample containing two images: $A_{\text{tibia}}\text{-HR}$ (high resolution image) and $A_{\text{tibia}}\text{-LR}$ (low resolution image). The samples A_{tibia} , B_{radius} , C_{radius} are considered as training dataset. D_{tibia} , E_{tibia} , F_{radius} , G_{radius} , H_{radius} serve as test set. For illustration, Fig. 5.2 presents 2D displays of samples A-D.

In the following, we address the selection of atoms and present 3D super resolution results.

5.3.1 Selection of 2D atoms

We have taken slices at interval equal to 10 along the axial direction, then we have cut slices into patches. In order to find a feature space as complete as possible but with less training data, a random sampling has been carried out, and finally 10,000 patch pairs have been considered as training input.

In the training output, we have observed that some atom pairs presented a common structural pattern, while some atoms were similar to noise. Since Φ^f and Φ^g are learned from the paired high and low resolution patches which display similar structures, and the update rule for the atoms are symmetric, d_i^f and d_i^g are expected to present similar structures, as indicated in (Yang et al., 2010a). Moreover, we have calculated the correlations between the atoms of the high and low resolution dictionaries with the correlation matrix $\Phi^{fT}\Phi^g$. In most cases, the correlations are maximum for the atoms corresponding to the same column index i and thus the same index corresponds to similar structure information, except for noise-like atoms. To quantify the coherence of each atom pair, we define the atom coherence measure as $M(i) = \langle d_i^f, d_i^g \rangle$, where $\langle \cdot, \cdot \rangle$ is the inner product. Then the atom pairs are arranged in a decreasing order of M , as shown in Fig.5.3 and Fig.5.4.

In Fig.5.3, we can observe that the earlier the atoms appear in the ordered dictionary Φ^f (with larger $M(i)$), the more obvious it is that they represent some structural information. Fig.5.4 presents the evolution of M arranged in decreasing order, $M(i+1)$ - $M(i)$ demonstrates the evolution speed of $M(i)$. The M slope shows a sharp change when $i = 250$, which suggests a transition from structure-like to noise-like atom pairs.

In order to confirm that $i = 250$ corresponds to a transition and to see the impact of these noise-like atoms, we use different amount of atom pairs to reconstruct the super resolution images. Fig.5.5 shows the evolution of PSNR and DICE as the number of atom pairs increases. It reveals that too many atom pairs in super-resolution stage degrade the reconstruction results, which may due to the introduction of noise-like atoms. On

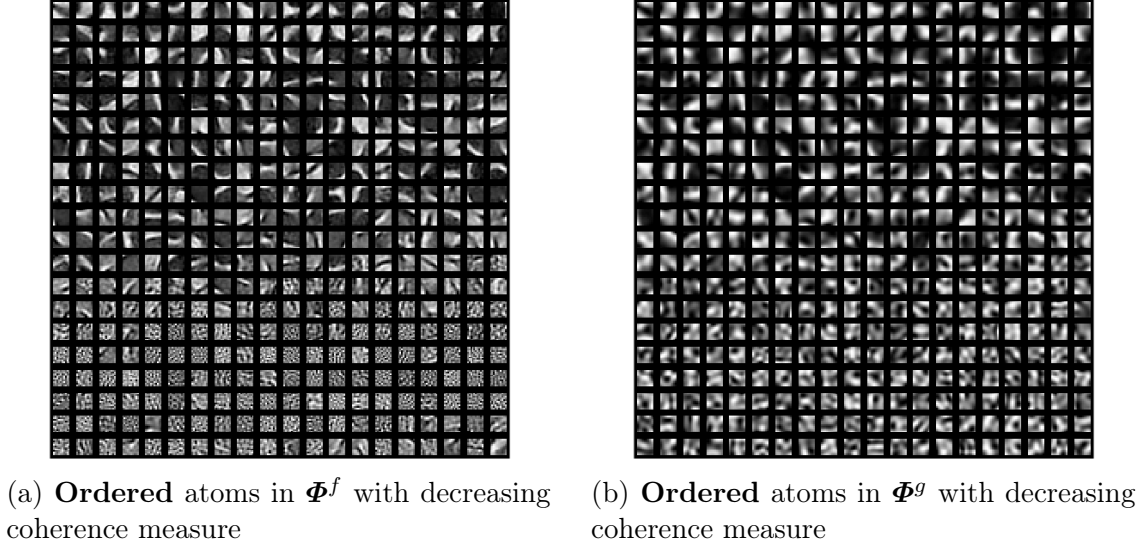


Figure 5.3: Illustration of atoms in sorted dictionaries, patch size is 10^2 , atom number equals to 400. Atoms are arranged in the decreasing order of $\mathbf{M}$.

the opposite, when the number of atoms pairs is too small, the results is not satisfying neither. We found that 250 atom pairs was a good choice for super resolution. This number coincides with the sharp decrease of $\mathbf{M}$ observed in Fig.5.4. We also performed tests with different patch sizes and over-completeness factors, similar conclusions have been drawn. Therefore, the coherence measure $\mathbf{M}$ proves to be a reliable tool to choose the number of atoms in super resolution reconstruction.

It is not possible to suppress noise-like atoms by retraining the dictionaries with less atoms. The noise on HR-pQCT images appears as an intrinsic information which can not be removed from primitive dictionaries during the training stage.

5.3.2 Results on 3D images

In this section, we present the results on 3D images obtained with the 2.5D strategy. One important hyper parameter is the size of patches. In this work, we varied the size of patches in the range $8^2, 10^2, 12^2, 14^2$ for 2.5D SCDL. Tab.5.1 and Tab.5.2 summarize the evolution of the different criteria as the patch size changes. Tab.5.1 presents the evolution of the different bone structure parameters for a training crop cut from A_{tibia} . It has 10% of slices appearing in the training set. Tab.5.2 describes the evolution of the criteria for a test volume whose slices are all out of training set. They both show that as the patch size increases, BV/TV tends to increase while Conn.D tends to decrease. Even though a large patch size leads to a better DICE, the BV/TV and Conn.D values are far from the reference value. As a result, we set patch size equal to 12^2 .

The 2.5 SCDL algorithm was applied to 5 test crops. The results are presented in Table 5.3, summarizing the bone parameters for the LR, SR (Super resolution) and HR images. All the parameters for the LR images are computed with raw data (without interpolation). The results show that, for most samples, both BV/TV and Conn.D are improved.

In 3D image reconstruction, we found that iterating 2.5D SCDL improved the results.

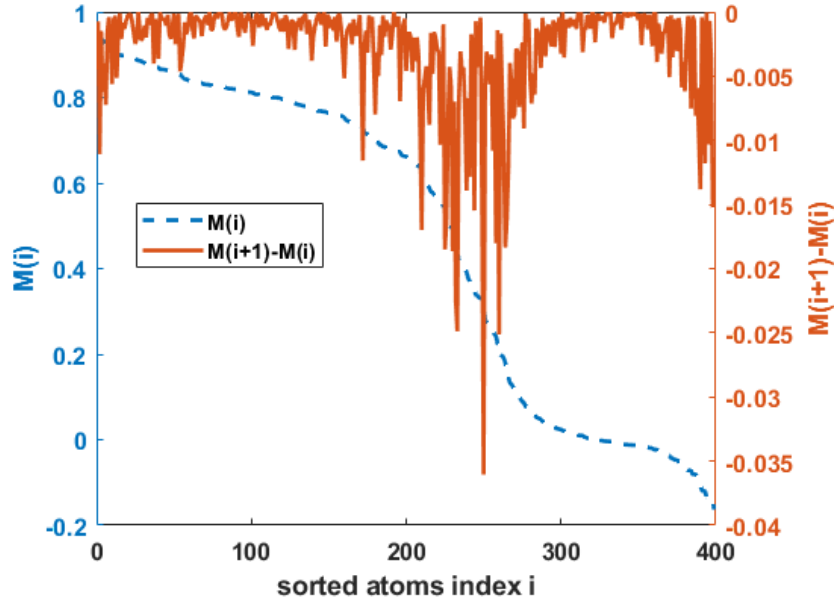


Figure 5.4: The plot of $\mathbf{M}$ (dashed line, blue) and of its forward derivative (plain, red) as a function of the sorted dictionary.

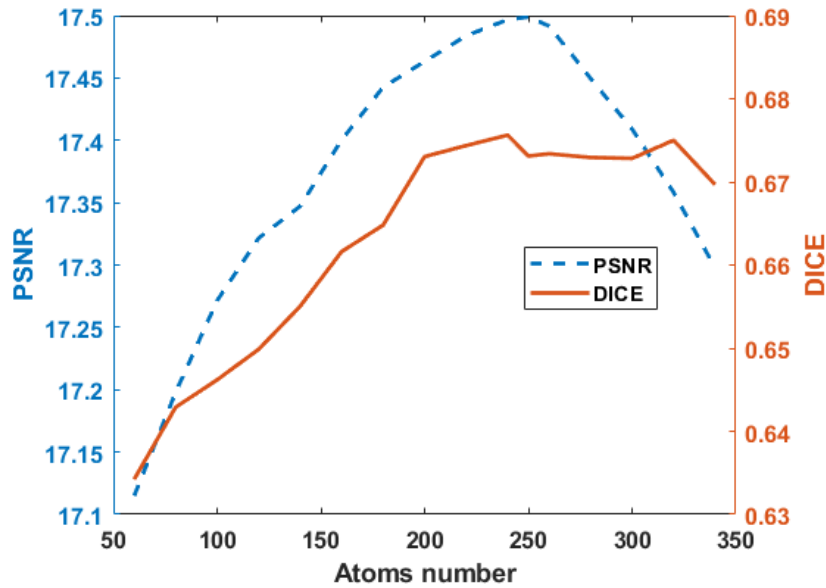


Figure 5.5: Evolution of two criteria, PSNR (dashed, blue) and DICE index (plain, red) when using different numbers of atom pairs.

Our experimental results show that a second iteration from (5) to (2) in Fig.5.1 can improve the image quality. This could be explained by the fact that some errors in (4.1)(4.2)(4.3) are corrected by taking the average. However, more iterations degrade the results because some details may be over corrected in (4.1)(4.2)(4.3). Similar to nonconvex inverse problems, an early stopping of the iterations can be interpreted as a regularization step.

Table 5.1: Image super-resolution criteria for the **training** crop cut from A_{tibia} obtained with different patch sizes. The over completeness factor is 4. 2.5D SCDL reconstructs 3D images with 2D dictionaries

patch size	PSNR	DICE	BV/TV(%)	Conn.D (mm ⁻³)
LR			40.1	5.17
HR (reference)		1	25.6	3.71
8 ²	15.66	0.771	29.4	4.59
10 ²	15.64	0.792	29.5	4.18
12 ²	15.91	0.797	29.7	3.82
14 ²	16.20	0.798	30.9	3.44

Table 5.2: Image super-resolution criteria for the **test** crop cut from D_{radius} obtained with different patch sizes. The over completeness factor is 4. 2.5D SCDL reconstructs 3D images with 2D dictionaries

patch size	PSNR	DICE	BV/TV(%)	Conn.D (mm ⁻³)
LR			33.10	6.51
HR (reference)		1	19.32	3.62
8 ²	15.02	0.752	22.60	3.60
10 ²	14.80	0.781	22.46	2.94
12 ²	15.31	0.791	22.36	2.81
14 ²	16.00	0.796	23.50	2.46

Table 5.3: Criteria for 5 test crops obtained with 2.5 SCDL.

parameter	D_{tibia}	E_{tibia}	F_{radius}	G_{radius}	H_{radius}
SR DICE	0.791	0.845	0.806	0.826	0.751
HR DICE	1.000	1.000	1.000	1.000	1.000
LR BV/TV(%)	33.10	29.13	29.43	33.33	36.03
SR BV/TV(%)	22.36	19.57	20.64	24.21	26.39
HR BV/TV(%)	19.32	18.92	18.52	20.70	24.07
LR Conn.D(mm ⁻³)	6.51	2.23	4.10	4.95	4.50
SR Conn.D(mm ⁻³)	2.81	1.32	2.57	3.24	3.77
HR Conn.D(mm ⁻³)	3.62	1.70	3.23	3.88	4.88

5.3.3 Comparison with TV based methods

In order to visually compare the different super resolution methods applied to the experimental CT images, Fig.5.6 and Fig.5.7 display 6 stacks of the sample A_{tibia} and H_{radius} cut from 3D volumes obtained with the different methods.

Compared with the high resolution images, TV regularization method could improve the quality of images, but they are still corrupted with noise. The image obtained with CH+TV seems to decrease the noise effects with respect to the TV images, but the

reconstructed micro-architectures are not so thin as the ones of the ground truth. 2D SCDL images are less noisy, but their structure is not smooth. 2.5D SCDL recovers a finer structure with a better texture.

Tab.5.4 and Tab.5.5 quantitatively summarize these differences. CH+TV, 2D SCDL and 2.5D SCDL improve the image quality compared to the low resolution image or to the TV regularization method. CH+TV is a post processing of TV methods, it indeed improves BV/TV, but tends to increase the Conn.D. Except for a better DICE, 2D SCDL has a similar performance as CH+TV, with a good BV/TV but much higher Conn.D. 2.5D SCDL has a better BV/TV and DICE than the other methods. Even though the Conn.D of the sample H_{radius} is not improved, we can conclude with the results in Tab.5.3, that 2.5D SCDL generally improves the Conn.D of samples. These conclusions need to be confirmed on larger data sets. The dictionary approach does not improve the PSNR. Nevertheless, the PSNR measures the similarity of gray level images and is less relevant to assess the quality of the segmented image. We have also performed a numerical spatial resolution analysis of the images for the different methods. The results showed that the spatial resolution is objectively improved with 2.5D SCDL and comparable to that of the HR image.

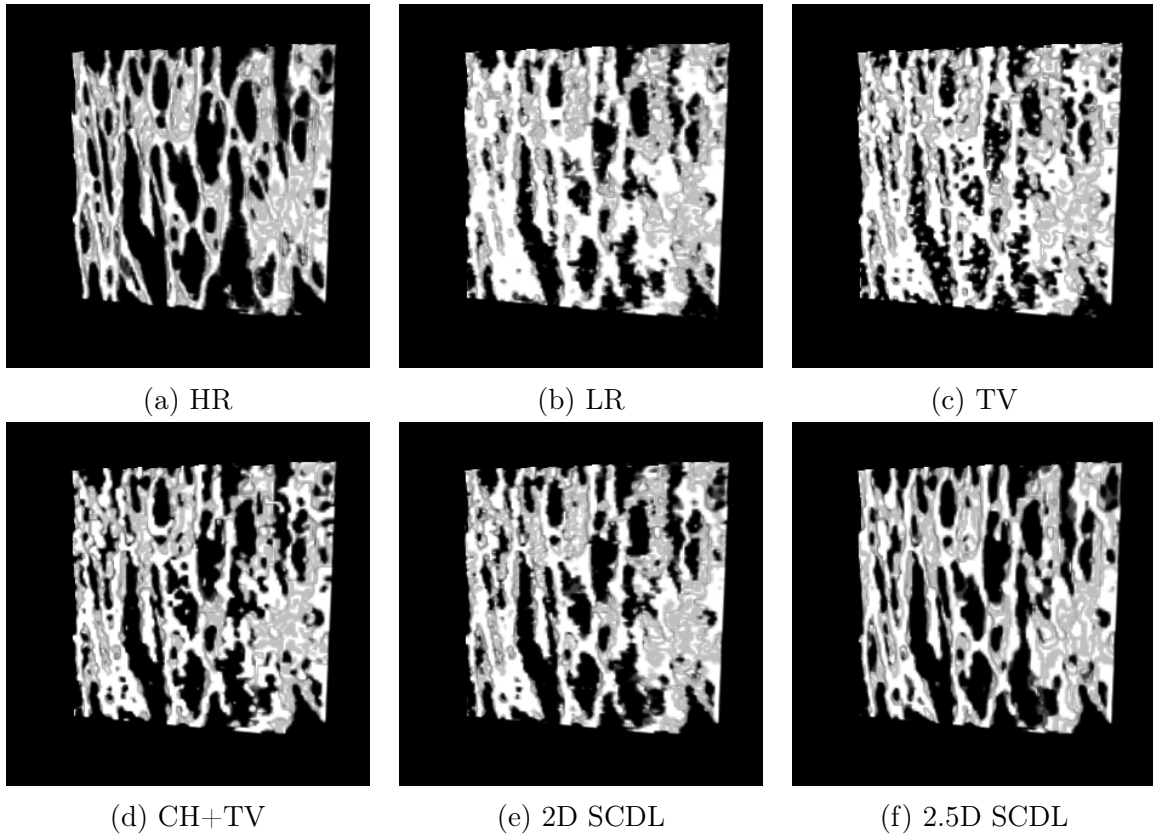


Figure 5.6: 3D rendering of images limited to a stack of 10 slices, on the first row, from left to right, high resolution, low resolution, TV super resolution; on the second row, CH+TV, 2D SCDL and 2.5D SCDL, respectively. The displayed image is a training crop cut from the sample A_{tibia} , the patch sizes of 2D and 2.5D SCDL are both 12^2 .

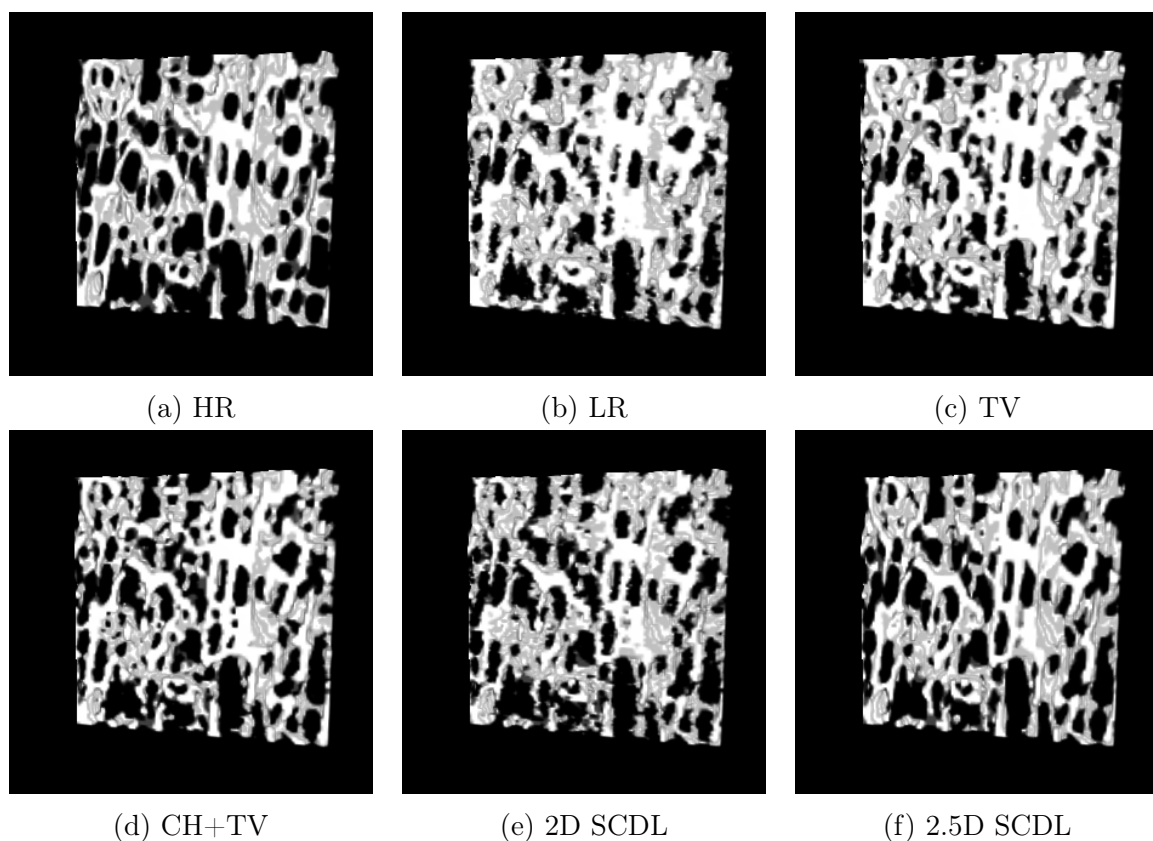


Figure 5.7: 3D rendering of images limited to a stack of 10 slices, from left to right, on the first row, high resolution, low resolution, TV super resolution, on the second row, CH+TV, 2D SCDL and 2.5D SCDL, respectively. Images is a test crop cut from the sample H_{radius} . The patch sizes of 2D and 2.5D SCDL are both 12^2 .

Table 5.4: criteria of the training crop cut from the sample A_{tibia} obtained with different super resolution methods. 2.5D SCDL reconstructs 3D images with 2D dictionaries

training crop	PSNR	DICE	BV/TV(%)	Conn.D (mm^{-3})
LR			40.1	5.17
HR (reference)		1	25.6	3.71
TV	18.87	0.742	39.0	4.25
CH+TV	17.53	0.779	29.0	4.44
2D SCDL	13.27	0.765	29.7	6.15
2.5D SCDL	15.91	0.797	29.7	3.82

5.4 Conclusion

In this paper, we developed a 2.5D SCDL method to solve the 3D HR-pQCT super resolution problem. The dictionary method has been implemented on slices in the three spatial directions.

The super resolution results are compared with the ones obtained with TV or TV combined with Cahn-Hilliard potential. The images restored by 2.5D SCDL are better

Table 5.5: criteria of the test crop cut from the sample H_{radius} obtained with different super resolution methods. 2.5D SCDL reconstructs 3D images with 2D dictionaries

test crop	PSNR	DICE	BV/TV(%)	Conn.D (mm^{-3})
LR			36.0	4.50
HR (reference)		1	24.1	4.88
TV	18.72	0.690	35.2	3.90
CH+TV	17.23	0.684	25.2	4.10
2D SCDL	13.59	0.707	26.3	4.72
2.5D SCDL	16.20	0.751	26.4	3.78

than the ones obtained with the other methods. We would like to apply this approach to larger experimental samples to validate the conclusion and extend the morphological analysis.

Even though the Conn.D of super resolution images is closer to the Conn.D of the ground truth, more research is expected to recover this bone micro-structure parameter at higher accuracy. In this work, we applied SCDL without considering the notion of cluster used in (Wang et al., 2012), because trabecular bone images may not have locality character as face images. But it's still a question for us whether or not locality constraints could improve the reconstruction results. Besides clusters, other constraints like low rank or l_1 regularization on the residual term (Jing et al., 2015, Wang et al., 2018) may potentially improve SCDL. The investigation of deep learning methods to solve this problem is also a very attractive perspective and will be considered in the next chapter.

5.5 Discussion

SCDL learns the dictionaries of high and low resolution images alternatively. It attempts to use two mapping matrices to linearly link both dictionaries. However, this fact may degrade the quality of super resolution reconstruction because the estimation of α^g depends on α^f .

Zeyde *et al.* in (Zeyde et al., 2010) proposed a method which get rid of such dependency. In the training stage, it first learns the dictionary of low resolution images, and determines the sparse coefficients α^g of low resolution patches. Afterwards, it estimates the high resolution dictionary with α^g . The low resolution images in the dataset is already interpolated, while the high resolution images in the dataset is the difference between the original high resolution images and the interpolated low resolution images. Alternatively, we could say that the high resolution patches in the dataset preserves the high frequency and details missing in the low resolution images. Several filters are designed to abstract low resolution features, and a Principle component analysis is applied to reduce the dimensionality of these features. Then K-SVD is applied to learn the dictionary of low resolution images. In the reconstruction stage, the super resolution image is estimated via the sparse coding vector of low resolution input image. This procedure keeps α^g from being affected by α^f , thus may improve the accuracy.

Inspired by (Zeyde et al., 2010), Gu *et al.* developed this idea and proposed convolution

sparse coding in (Gu et al., 2015). They also introduce a mapping matrix, which permits to use shorter sparse coding vector to decompose low resolution images and longer sparse coding vector for high resolution images. Compared to (Zeyde et al., 2010), another improvement of (Gu et al., 2015) is that it can directly deal with low resolution images, thus reducing the computation burden.

Analysis operator learning approaches are similar to dictionary learning methods. Considering the problem to recover an image $\mathbf{f}$ based on the observation $\mathbf{f}_{obs}$, it is formulated as follows in dictionary learning:

$$\begin{aligned} \arg \min_{\alpha^f} & \|\mathbf{f}_{obs} - \Phi^f \alpha^f\|_2^2 + \mu \mathcal{R}(\alpha^f) \\ \mathbf{f} &= \Phi^f \alpha^f \end{aligned} \quad (5.10)$$

where $\mathcal{R}(\cdot)$ serves as a regularization term enhancing the sparsity. In analysis operator learning, the problem is written as:

$$\arg \min_{\mathbf{f}} \|\mathbf{f} - \mathbf{f}_{obs}\|_2^2 + \mu \mathcal{R}(\Psi \mathbf{f}) \quad (5.11)$$

When the dictionary Φ is invertible, the analysis operator Ψ can be written as $\Psi = \Phi^{-1}$, both approaches become equivalent. The regularization term in Eq. 5.11 can be parametrized. Let us assume it depends on some parameters β . The former equation can be rewritten:

$$\arg \min_{\mathbf{f}} \|\mathbf{f}(\beta) - \mathbf{f}_{obs}\|_2^2 \quad (5.12)$$

with $\mathbf{f}(\beta) = \arg \min \mathcal{R}(\Psi(\beta) \mathbf{f})$. The first problem is the upper optimization problem and the second one is the lower optimization one. Chen *et al.* in (Chen et al., 2014) use bi-level optimization to solve this problem. They use the lower level problem to optimize latent image, and the upper level problem to choose the parameters characterizing the regularization terms and the analysis operator.

In this chapter, all the atoms in the dictionaries are in 2D. We agree that that atoms in 3D may abstract structural information more accurately than 2D. However, the training time is not acceptable because the atoms are learned sequentially. Moreover, taking the choice of parameter into account, it takes even longer time to use the method. Furthermore, the mapping matrix $\mathbf{W}$ describes a linear mapping function, which is not sufficient for complex mapping functions. For these reasons, we did not continue our work on 3D dictionary learning and attempt to use deep learning framework to solve the our super resolution problem. The relevant work is presented in the following chapters.

Chapter 6

Deep learning for super resolution

In previous works, we have applied variational methods such as TV based regularization for our super resolution problem (Toma et al., 2014b, Peyrin et al., 2015, Li et al., 2017a). The results showed that TV smooths images, but the images suffered from staircase noise artifacts. By adding a double-well regularization to TV, the contrast of the images was enhanced but the curvature of the trabecular shapes was degraded, as explained in chapter 4. Chapter 5 presents an approach inspired from the semi-coupled dictionary learning (Wang et al., 2012) with the purpose of recovering trabecular micro-architecture details. The resolved images have been improved quantitatively and visually. In fact, one challenge to use dictionary learning methods for super resolution is to couple efficiently the high and low resolution dictionaries. Different mapping functions have been discussed in (Yang et al., 2010a, Yang et al., 2012, Wang et al., 2012, He et al., 2013). Yet, these mapping functions are assumed to be linear, which may be not sufficient for complex mapping functions. Deep learning network is a good candidate to describe nonlinear mapping functions.

Before introducing the application of deep learning network in our super resolution problem, we would like to present a review of the application of deep learning network in super resolution problems.

6.1 The state of the art on deep learning for super resolution

In this section, we first demonstrate different deep learning frameworks for super resolution problems, then compare their performances.

6.1.1 Deep learning network for super resolution problems

Dong *et al.* applied CNN to solve super resolution problems (Dong et al., 2014) (SRCNN). They constructed a 3 layers CNN, as illustrated in Fig.6.1. The first layer could be considered as sparse coding of low resolution image features, the second layer carried on a non-linear mapping to get sparse coefficients for high resolution image features, and the third layer reconstructs super resolution images. The approach is similar to dictionary learning approach. But it adds nonlinearity and could be extended to a deeper architecture

to improve the accuracy of representation.

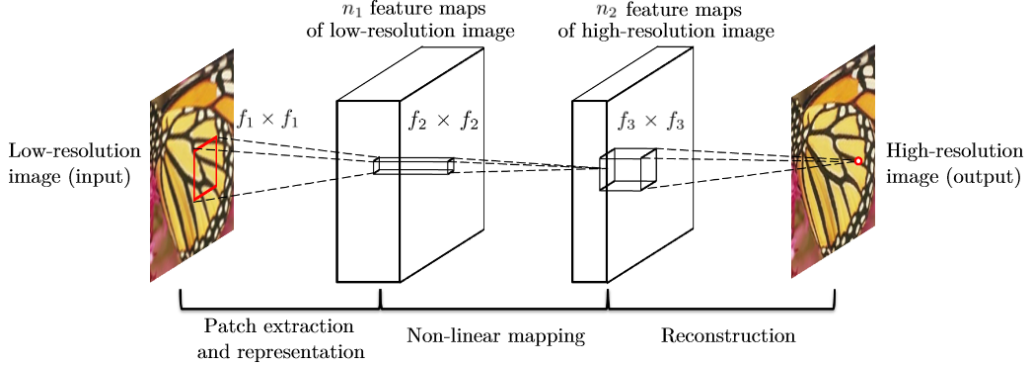


Figure 6.1: Illustration of SRCNN architecture, this figure is collected from (Dong et al., 2014)

SRCNN (Dong et al., 2014) is a cornerstone in the literature of deep learning for super resolution problems. It seldom disappears in the benchmarks of CNN based approaches. Nevertheless, the SRCNN network is usually blamed for the short network depth.

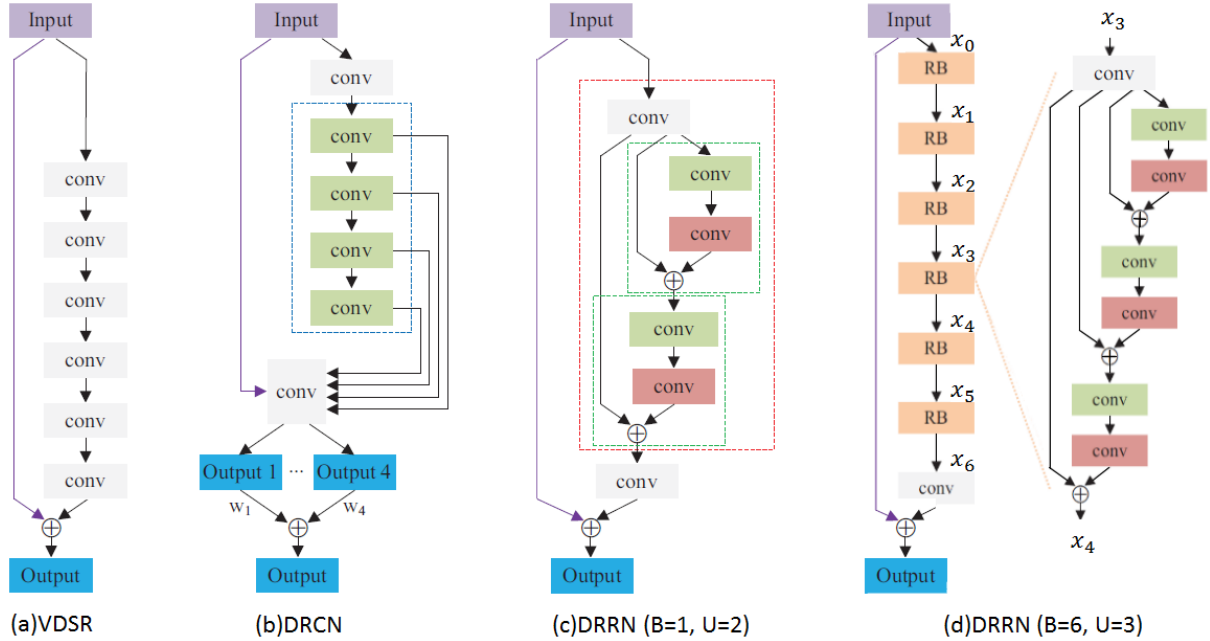


Figure 6.2: Figure adapted from (Tai et al., 2017). (a) VDSR (Kim et al., 2016a). The skip connection between the input and output ensures the stability of the network, particularly for deep network. (b) DRCN (Kim et al., 2016b). The dashed green block is a recursive block. The convolution layers within the recursive block share the same parameters. (c) DRRN (B=1, U=2) (Tai et al., 2017). B denotes the number of recursive blocks and U is the number of recursive units in a recursive block. The green dashed block denotes a residual unit, the red dashed block represents a recursive block. Inside of a recursive block, the convolution layers in the same color share the same parameters. (d) DRRN (B=6, U=3) (Tai et al., 2017). The orange 'RB' block denotes a recursive block.

Kim *et al.* later proposed a Very Deep residual network for Super Resolution (VDSR) (Kim *et al.*, 2016a). VDSR has very deep architecture (20 layers) and each layer consists of small filters. The skip connection from the input image to the output estimation, as shown in Fig.6.2(a), forces the convolution filters to learn the residual between the estimation and the ground truth images. This is also the reason why it is named as residual network. The gradient clipping strategy allows to train the network with high learning rate, thus accelerates the convergence speed despite the huge size of the architecture. The principle of gradient clipping is to truncate the individual gradient so that all the gradients are constrained in a predefined range (Kim *et al.*, 2016a). The authors found that increasing the depth of the networks improves the accuracy of the results.

Deeply Recursive Convolutional Network (DRCN) (Kim *et al.*, 2016b) uses a recursive structure so that the length of the network is increased while the number of parameters is reduced. The recursive structure uses the same simple filters repetitively to abstract image features. As shown in Fig.6.2(b), the green dashed rectangular is a recursive block. All the convolutions marked in green within this recursive block share the same parameters. All the intermediate outputs from recursive block and the input of the network are then fed into a convolution layer to generate output predictions. In Fig.6.2(b), there are 4 output predictions. The final estimation is determined by a linear combination of the output predictions, and it is optimized by a squared mean loss.

One limit to the performance of general recursive networks is that the gradient can explode or vanish, which induces instability and reduces the learning ability of the network. Since the network is optimized with back propagation, all the parameters are updated with the gradient chain. With the multiplicative rule of gradient chain, the gradient of parameters may explode or vanish. The authors of (Kim *et al.*, 2016b) tackled this problem with two strategies: recursive supervision and skip connection. Recursive supervision means that all the intermediate outputs from recursive block participate in the determination of output predictions, and each output prediction is supervised by a mean squared loss. The differences between the output predictions smooth the gradient of parameters. Moreover, the skip connection between the input of the network and the outputs of the recursive block makes that the network needs less recursion layers, thus it alleviates the gradient explosion and vanishing problem, according to (Kim *et al.*, 2016b). In my opinion, the recursive supervision is a remedy for vanishing gradient, and the skip connection avoids gradient exploding. These ideas are similar to the ones on which VDSR is based. DRCN indeed reduces the number of parameters, nevertheless, the memory to save intermediate outputs can not be ignored.

Similarly to the DRCN, the Deep Recursive Residual Network (DRRN) (Tai *et al.*, 2017) applies recursive learning. But contrary to DRCN, the recursive unit in DRRN is a modified resnet unit, as shown in Fig.6.2(c). The green dashed block denotes a modified resnet which consists of two convolution layers, and each convolution layer is a stack of batch normalization, ReLU activation function followed by a weight layer (convolution filters). The batch normalization outputs batches with zero mean value and standard deviation equal to 1. It helps to increase the learning rate and makes the network more robust to the initialization (Ioffe and Szegedy, 2015). The red dashed block is the recursive block of the network. Within the recursive block, the convolution layers marked in the same colors share the same parameters. The skip connections in the architecture avoids the problem of gradient vanishing and explosion, moreover, it allows the network to learn

complex functions. DRRN can be parametrized by B and U , where B is the number of the recursive blocks, U denotes the number of the recursive units within one recursive block. In Fig.6.2(c), the network contains 1 recursive block, and this recursive block has 2 residual units ($B = 1, U = 2$). Fig.6.2(d) is composed of 6 recursive blocks and every block has 3 residual units ($B = 6, U = 3$). In fact, the increase of recursive block quantity increases the number of parameters in the network, while the increase of residual unit quantity does not change the amount of parameters but increases the depth of the network. The authors in (Tai et al., 2017) noted that even though the DRRN networks are parameterized by B and U , the networks' efficiencies are comparable if their depths are similar.

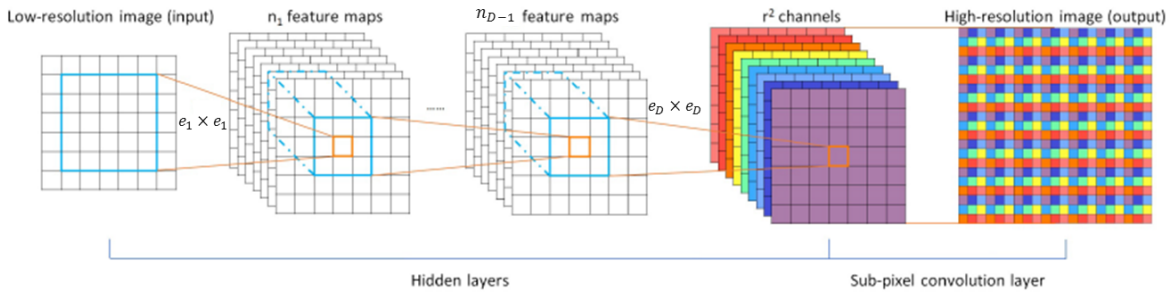


Figure 6.3: Illustration of ESPCN architecture. This figure is collected from (Shi et al., 2016)

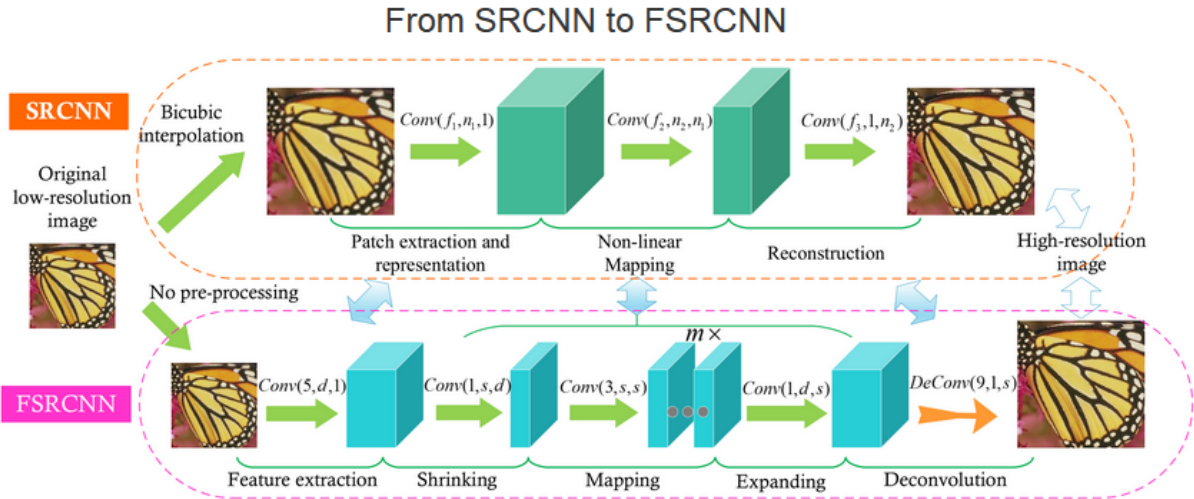


Figure 6.4: Illustration of FSRCNN architecture. This figure is collected from (Dong et al., 2016).

All the three architectures presented in Fig.6.2 use interpolated low resolution images as input. This operation increases the burden of the calculation and slows down the efficiency of the network. Shi *et al.* in (Shi et al., 2016) proposed Efficient Sub-Pixel Convolutional neural Network(ESPCN) where the input is the original low resolution images. ESPCN proceeds upsampling operation at the sub-pixel convolution layer, as shown in Fig.6.3. The stride of sub-pixel convolution is $1/r$, where r is the up-sampling

factor. The stride can be interpreted as the step size of convolution. The sub-pixel convolution increases features scale by rearranging the pixels learned in the former layer. As illustrated in Fig.6.3, the input channels of the sub-pixel convolution layer are marked in different colors and their correspondent pixels are regularly distributed after sub-pixel convolution.

Another method of upsampling operator is proposed in the Accelerated SRCNN (Dong et al., 2016)(FSRCNN), which is an extension of SRCNN. Compared with SRCNN, the non-linear mapping function in FSRCNN is more flexible and robust. Moreover, two additional layers have been introduced to reduce the number of parameters while keeping the performance of the network. A brief version of FSRCNN achieves real-time processing speed. Differing from ESPCN, FSRCNN applies deconvolution layer with stride of r for upsampling operation. It's noteworthy that the "deconvolution" here is a transpose convolution. In image segmentation, deconvolution has another meaning, it can be regarded as an inverse operation of pulling (Zeiler and Fergus, 2014).

ESPCN (Shi et al., 2016) and FSRCNN (Dong et al., 2016) are two pioneer works which integrate the upsampling operation within the network so that no interpolation is necessary. This is a big progress for the development of deep learning in super resolution problems. Researchers extended these two ideas in their work (Lim et al., 2017, Lai et al., 2017, Zhang et al., 2018) to further improve the efficiency of deep learning based methods for super resolution problems, different upsampling factors have been investigated and different methods have been tested on the same datasets. The Set5(Bevilacqua et al., 2012) and set14(Zeyde et al., 2010) are two commonly used test sets for super resolution benchmarks, which include 5 and 14 images respectively. Fig.6.5 illustrates 8 test images among which 4 images are from the Set5, the other 4 are in the Set14.

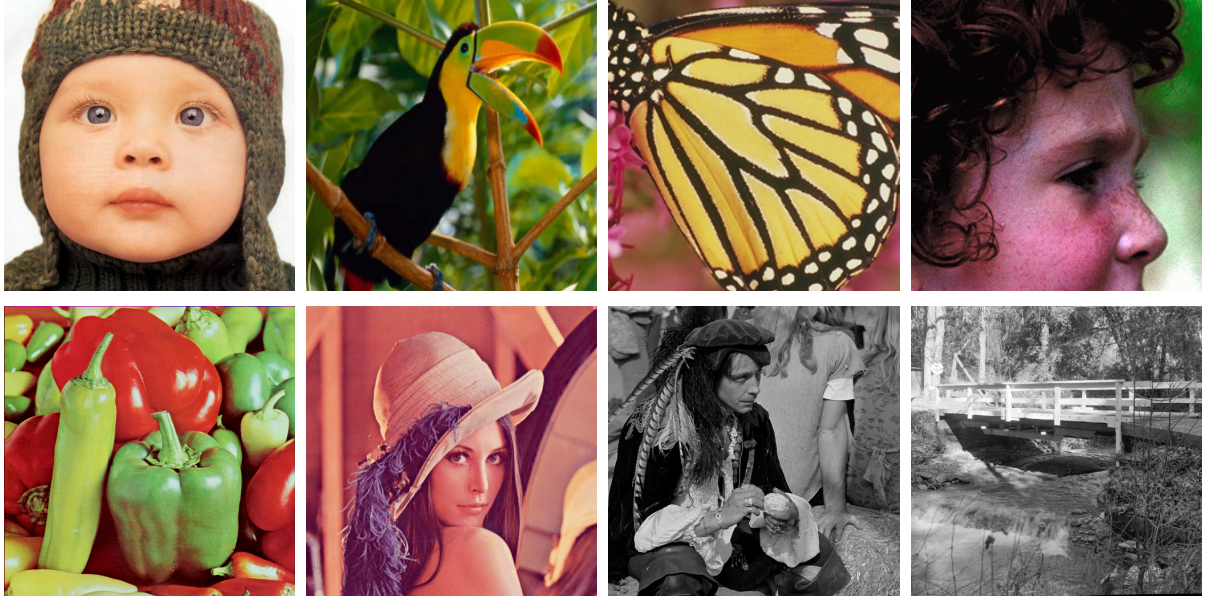


Figure 6.5: Illustration of 8 images in the set5 as well as set14. The images on the top row are 4 images from the Set5. The images on the bottom row are 4 images from the Set14.

The Enhanced Deep Residual Networks (EDSR) (Lim et al., 2017) aims to use residual

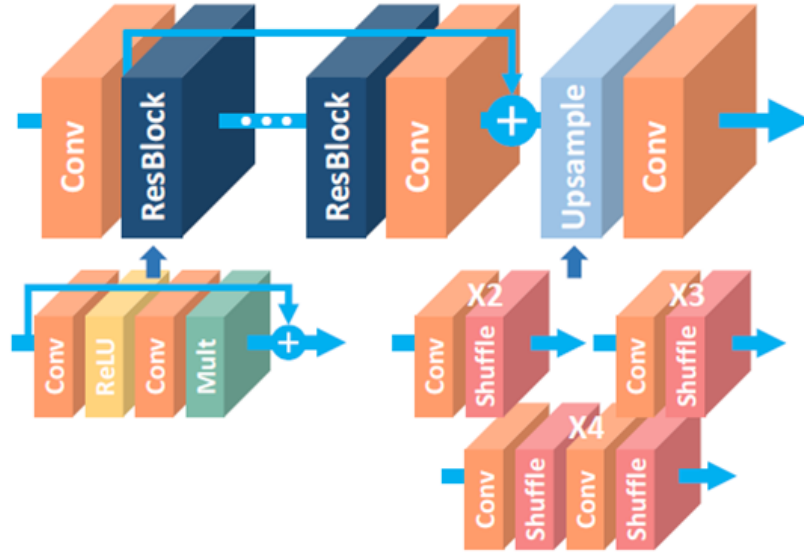


Figure 6.6: Illustration of EDSR network, the plot is collected from (Lim et al., 2017). The ResBlock is stacked by a convolution layer, Relu activation, convolution layer and a constant scaling layer. The constant scaling layer simply rescales the residual with the purpose of keeping the stability of the network when the number of filters exceeded 1000 (Szegedy et al., 2017).

block to enhance the structural information at low resolution scale, then upscale images at the last second layer. The structure of EDRN is drawn in Fig.6.6. In the 'ResBlock', the authors removed the batch normalization and introduced a constant scaling layer (the green layer marked with 'Mult' in Fig.6.6). They think the suppression of batch normalization reduces memory consumption and keeps the range flexibility of features. A constant scaling layer is proposed in (Szegedy et al., 2017). Szegedy *et al.* mentioned that when the number of filters exceeds 1000, the residual variants become instable and tend to lose their learning ability. Neither batch normalization nor decreasing the learning rate can solve this problem. Nevertheless, rescaling the residual with factor 0.1 before adding to the identity mapping seems to be able to keep the network stable. For this reason, there is a constant scaling layer in 'ResBlock'. The upsample operation is sub-pixel convolution, following the idea proposed in (Shi et al., 2016).

Lai *et al.* proposed a deep Laplacian Pyramid Networks (LapSRN) in (Lai et al., 2017) for super resolution problems. The main idea is to gradually upscale features. Its architecture has two branches: one for feature extraction, the other serves for reconstruction, as presented in Fig.6.7. In feature extraction branch, convolution layers marked in red abstract features' characters, deconvolution (transposed convolution) layers marked in blue upscale features. The output of deconvolution layers connects with 2 layers: one serves for residual information in the image reconstruction branch; the other is used for feature extraction for the next upsampling operation. The deconvolution layers in image reconstruction branch are initialized with bilinear kernel which is essential to force the feature extraction branch to learn residual features.

It seems that the image reconstruction branch is responsible to learn low frequency information, and the feature extraction branch refines the details and feeds high frequency

information to the image reconstruction branch.

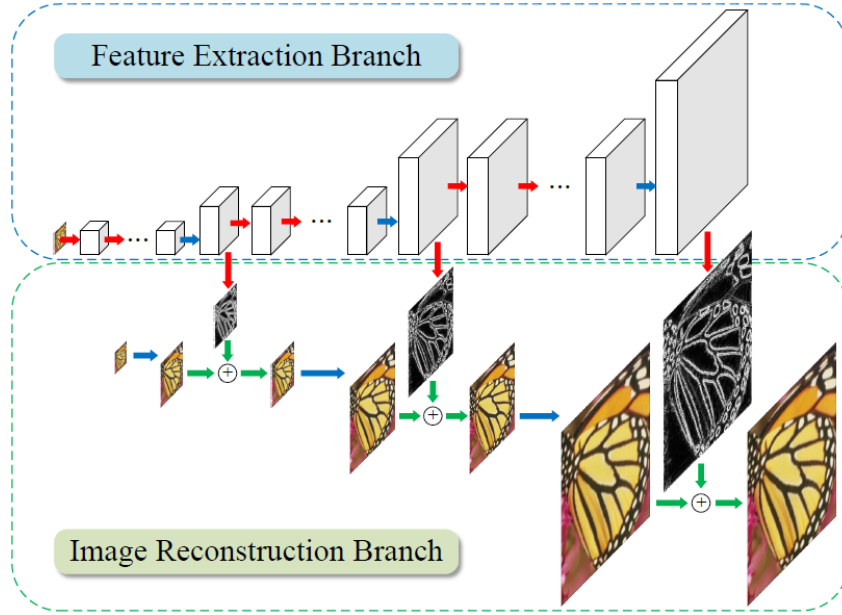


Figure 6.7: Illustration of LapSR, this plot is collected from (Lai et al., 2017). Red arrows refers to convolution layers, blue arrows are transposed convolutions, green arrows are element-wise addition.

Ledig *et al.* incorporated VGG network (Simonyan and Zisserman, 2014) within Generative Adversarial Network (GAN) (Goodfellow et al., 2014) to deal with super resolution problems (Ledig et al., 2016) (SRGAN), as shown in Fig.6.8.

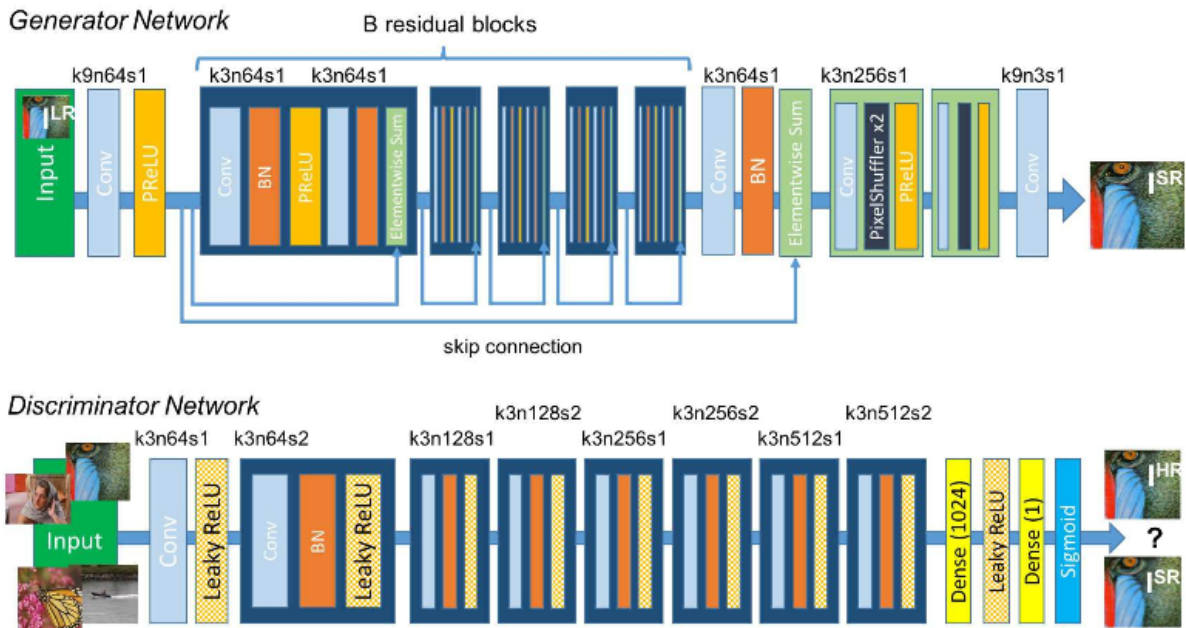


Figure 6.8: Illustration of SRGAN, the diagram is collected from (Ledig et al., 2016). k represents the size of filter, n as the number of feature maps, s as the stride for each convolutional layer.

VGG is a very deep CNN, it uses small filters to capture image features at different scale. With the increase of the layer's depth, the size of features decreases and the number of channels increases. Compared with mean square error distance, VGG network is less sensible to the changes at pixel level. GAN is composed of a generator and a discriminator. The generator is used to output estimation, and the discriminator constrains the estimation to follow some predefined distribution. The authors in (Ledig et al., 2016) use generator to generate super resolution image, and the discriminator is used to distinguish the super resolution image from the high resolution image. If the discriminator can not see the difference between the estimated super resolution image and the high resolution image, this estimation is assumed to be a good approximation of the high resolution image. Precisely, the generator enhances the similarity of the estimation and the ground truth in the space described by VGG network. The discriminator diminishes the difference between the estimation and the ground truth.

It can be seen that the skip connection has a significant impact on deep neural network, and the residual network or resnet unit plays an important role in the recent advanced deep learning network designed to solve super resolution problems. Huang *et al.* introduced the densely connected CNN (denseNet) in (Huang et al., 2017) to solve object recognition, where more skip connections have been introduced compared with residual network or ResNets. Fig.6.9 demonstrates the framework of the denseNet. Differing from ResNets, the denseNet concatenates the outputs of former blocks while ResNets uses summation. The denseNet encourages to reuse the features, as a result, less parameters are required.

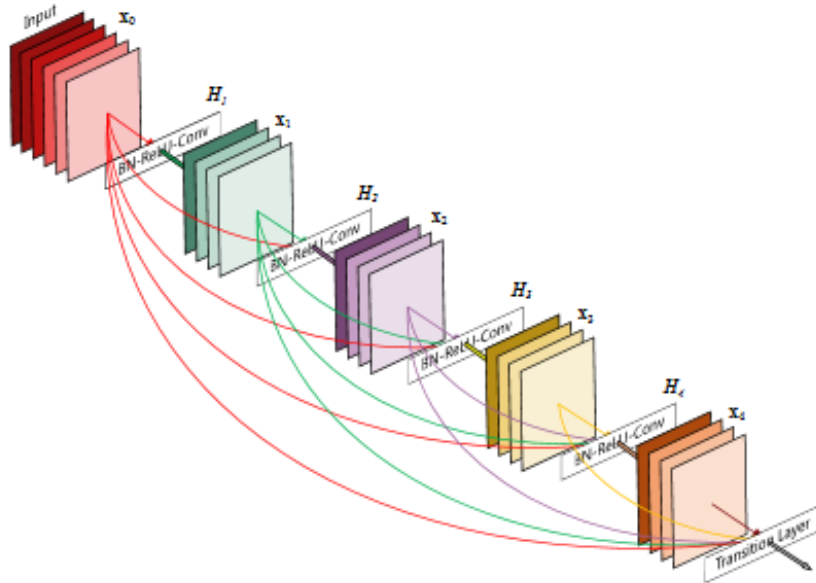


Figure 6.9: Illustration of densely connected CNN, this figure is collected from (Huang et al., 2017)(denseNet).

The success of denseNet quickly draws the attention of researchers in super resolution field (Tong et al., 2017, Zhang et al., 2018).

Tong *et al.* in (Tong et al., 2017) build an framework for super resolution problems using dense skip connections (SRdenseNet). The schema is illustrated in Fig.6.10. As

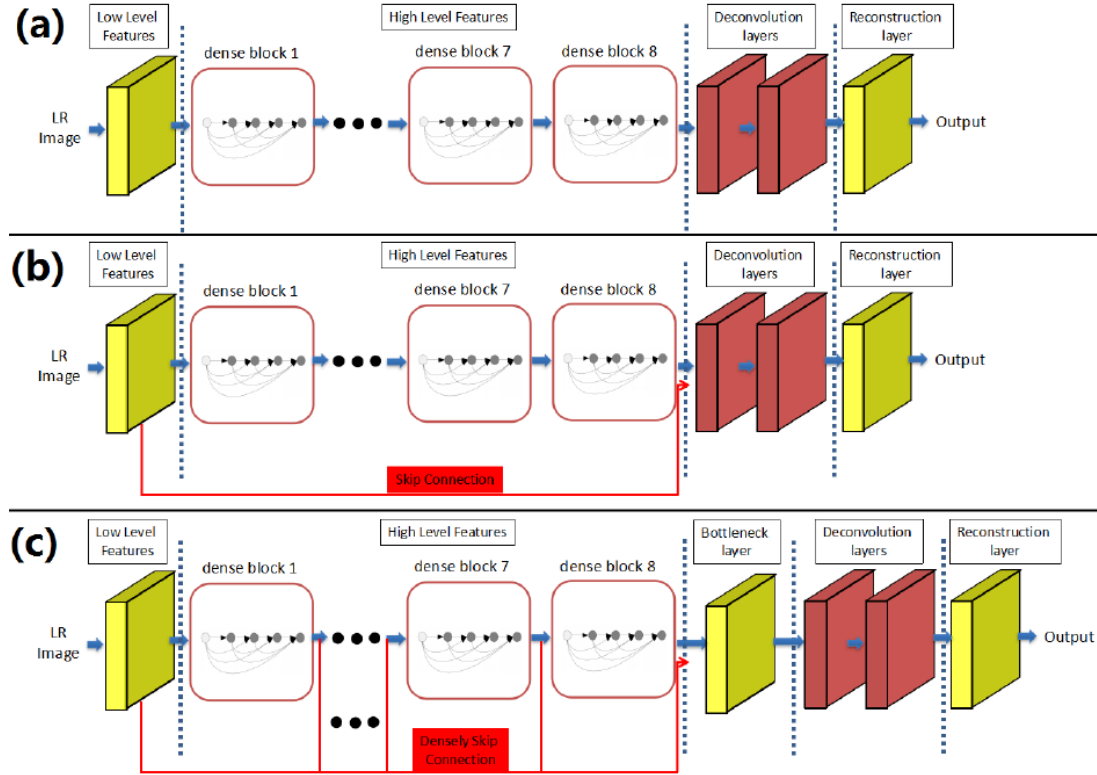


Figure 6.10: Illustration of the architecture of SRdenseNet, this diagram is collected from (Tong et al., 2017). (a) uses dense block to enhance image features then feed the enhanced features into deconvolution layers for upsampling. (b) long skip connection preserves low level features and enables the dense blocks to learn high level features. (c) The bottleneck combines both low and high level features and feed the results into deconvolution layers.

explained in the caption of Fig.6.10, the SRdenseNet extends the conception of denseNet to blocks for the purpose of abstracting high level features, and combines all the high level features with low level features via bottleneck layer, afterward passing through deconvolution layers for upsampling. The bottleneck reduces the dimensionality of the features. The way that the SRdenseNet combines low and high level of features is similar to the one in DRCN (Fig.6.2(b)), but differing in the reuse of parameters.

Zhang *et al.* proposed a residual dense network to solve super resolution problem in (Zhang et al., 2018)(RDN). Fig.6.11 compares different network blocks: residual block, dense block, residual dense block. As can be seen in the Fig.6.11, residual dense block combines residual block and dense block. The 1×1 convolution layer is used to reduce the dimensionality. The residual method forces the filters to learn residual part information. For instance, in VDSR, the long skip connection conveys low frequency information to the output so that the convolution layers in the network are forced to learn high frequency information. Therefore, the learning task is simplified. The dense block boosts the ability of the network to describe complex functions. The residual dense block takes in the advantages of both Residual block and Dense block, thus is expected to give a better performance.

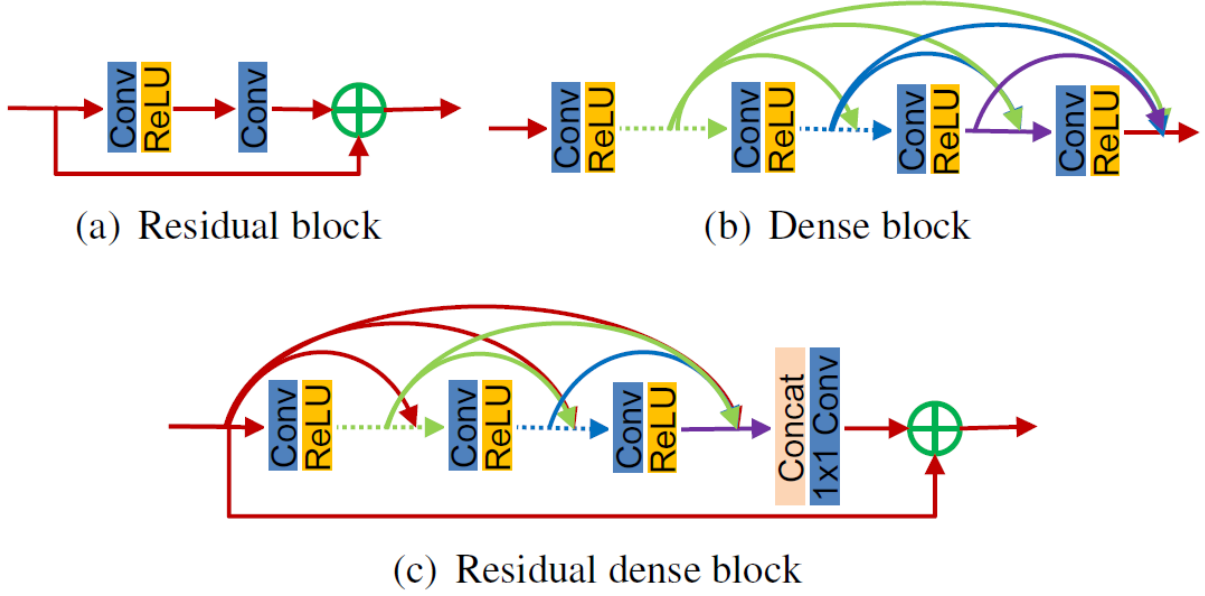


Figure 6.11: Comparison of different network blocks, this image is collected from (Zhang et al., 2018). (a) Residual block in (Lim et al., 2017). (b) Dense block in SRDenseNet (Tong et al., 2017). (c) Residual dense block (Zhang et al., 2018)

Moreover, if we say the residual dense block is a micro network structure, Zhang *et al.* in (Zhang et al., 2018) projects this structure into the macro scale. Fig.6.12 shows the entire architecture of RDN where the residual dense structure is flexibly used.

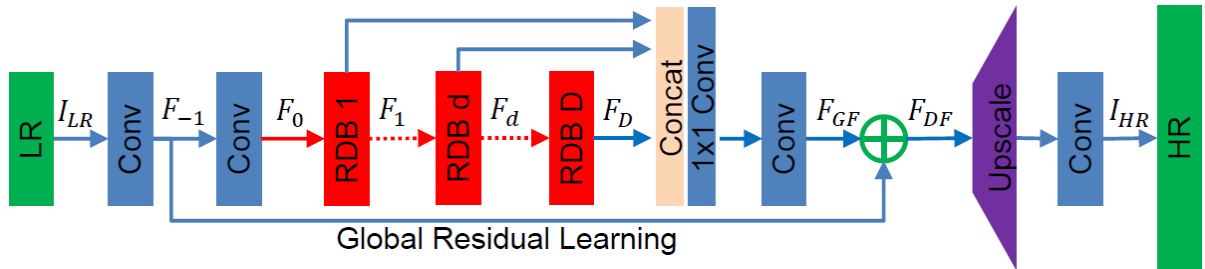


Figure 6.12: The illustration of the residual dense network, the plot is collected from (Zhang et al., 2018).

There are three major differences between SRDenseNet and RDN. First, in the scale of block unit: RDN has a long skip connection similar to residual block which SRDenseNet does not have. Second, the connections between block units: SRDenseNet does not have connections among the blocks while in RDN, the blocks are densely connected, following the architecture of denseNet. Third, in the scale of global structure: in SRDenseNet, the low and high level features are concatenated together before passing through the bottleneck; but in RDN, only the high level features are densely connected, then combined with low level features via residual structure.

6.1.2 Comparison of the performances of the different networks

Tab.6.1 compares the performance of different architectures in terms of PSNR. The test set is two common natural image sets: set 5 and set 14. The RDN+ outperforms all the other methods in the table. The sign '+' denotes that the results are improved with self-similarity (Lim et al., 2017). As we know, during the training stage, data augmentation increases the amount of training data, thus boosts the network. Self-similarity has the similar idea, but it is employed during the reconstruction stage. It assumes that the network is invariable to the geometrical transformation of the input data. For example, a test sample has 8 augmented inputs after geometric transformation (including identity). These 8 samples will be fed into the network and generate 8 estimations. The final result is the average of these 8 estimations after the correspondent inverse geometric transformation.

Table 6.1: Comparison of PSNR different methods for natural image sets set5 and set14, the low resolution images are generated with BI degradation model. The upsampling factor is 4. The sign '+' represents that the approach has been boosted with self-assemble. The results are collected from correspondent publications.

methods	set5	set14
SRCNN(Dong et al., 2014)	30.49	27.61
ESPCN(Shi et al., 2016)	30.90	27.73
FSRCNN(Dong et al., 2016)	30.55	27.50
LapSRN(Lai et al., 2017)	31.33	28.06
VDSN (Kim et al., 2016a)	31.35	28.03
DRCN(Kim et al., 2016b)	31.53	38.04
DRRN(B1U25)(Tai et al., 2017)	31.68	28.21
EDSR(Lim et al., 2017)	32.46	28.80
EDSR+(Lim et al., 2017)	32.62	28.94
SRGAN(Ledig et al., 2016)	29.40	26.02
SRDenseNet(Tong et al., 2017)	32,02	28,50
RDN(Zhang et al., 2018)	32.47	28.81
RDN+(Zhang et al., 2018)	32.61	28.92

In this chapter, we introduce some deep learning networks for super resolution problems. Evolving from the SRCNN to deep network such as VDSR, from recursive network to recursive residual network, from residual network to densely connected network, the deep learning architecture tends to be more and more robust and more densely connected.

6.2 The application of deep learning methods in medical images super resolution problems

In medical image processing, the resolution of in vivo image scanning is limited because of the acquisition time or the quantity of dose. In this section, we are going to review the application of deep learning methods in medical images super resolution problems. The

literature is quite recent, yet there is already a strong interest on this topic of applications for various imaging modalities.

Umehara *et al.* investigate the application of SRCNN for enhancing image resolution on chest CT images (Umehara *et al.*, 2017). They show that SRCNN outperforms traditional linear interpolation methods.

Mansoor *et al.* investigates the application of SRGAN (Ledig *et al.*, 2016) for two imaging modalities (CT and MRI) in (Mansoor *et al.*, 2018). They applied VGG-like network to abstract the features, thus avoiding to enhance the similarity based on pixel-level. The basic mechanism is similar to the SRGAN introduced in the last section. The network is trained with 2D slices in three plane. This research work would be even more interesting if a comparison with 2D and 3D was given. In their training set, the low resolution images are generated by down sampling the high resolution images smoothed with Gaussian.

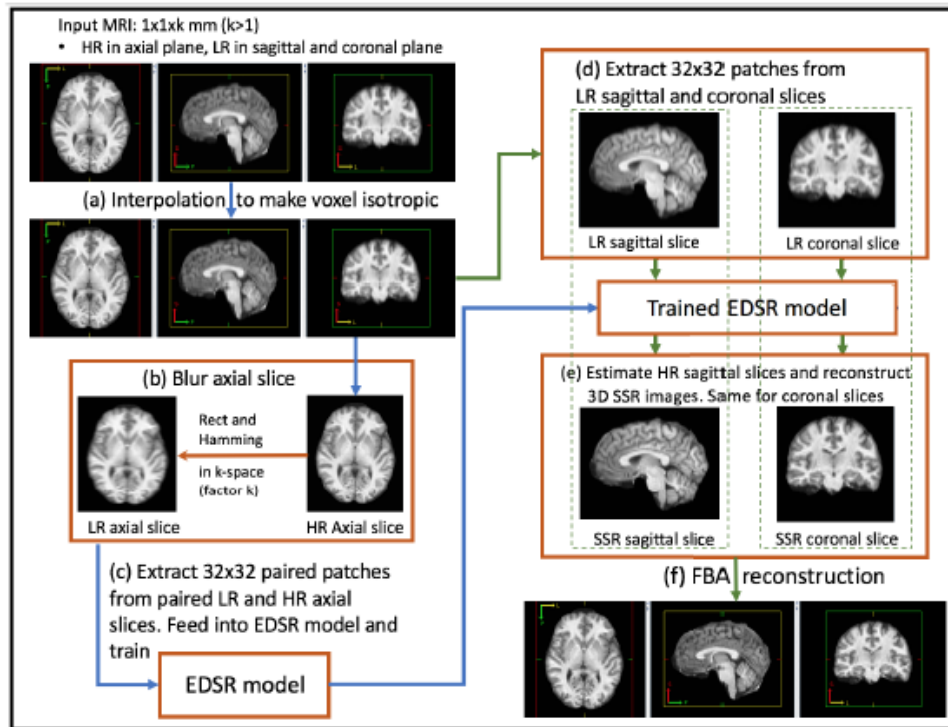


Figure 6.13: The proposed framework for MRI super resolution, this figure is collected from (Zhao *et al.*, 2018).

Zhao *et al.* extends EDSR (Lim *et al.*, 2017) in MRI super resolution reconstruction (Zhao *et al.*, 2018). Their proposed network is illustrated in Fig.6.13. In their original database, the images in axial plane is at high resolution, but the sagittal and coronal plan are at low resolution scale. They propose to artificially degrade the high resolution image in axial plan to generate low resolution images in axial plane. Afterward, they train EDSR with the paired low and high resolution images which are in axial plane. Then the low resolution image in sagittal and coronal plane will be reconstructed via the trained EDSR model. This approach assumes that the degradation kernel is isotropic in three dimensions, which is crucial for the quality of reconstruction.

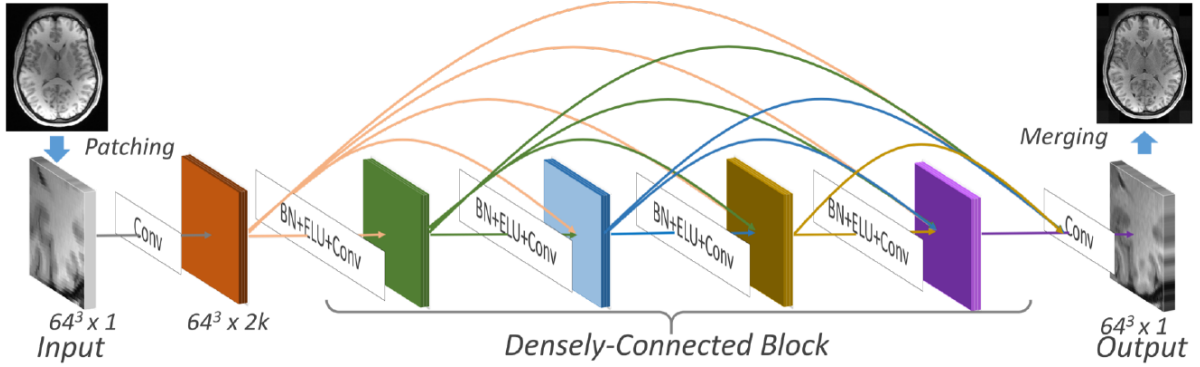


Figure 6.14: Framework of 3D DCSRNet, this plot is collected from (Chen et al., 2018).

Chen *et al.* in (Chen et al., 2018) proposed a densely connected super resolution network (DCSRNet) for brain MRI images. The output of each layer is connected to the following layers. They conclude that the neural network for 3D is better than their 2D counterparts, the proposed method outperforms 3D FSRCNN. Yet, very few details are given. The proposed structure is demonstrated in Fig.6.14. It can be seen that the network is densely connected, the authors in (Chen et al., 2018) resumed 3 advantages of their proposed network: faster training since the path is shorten; tiny model because of weight sharing; less overfitting due to the reuse of features.

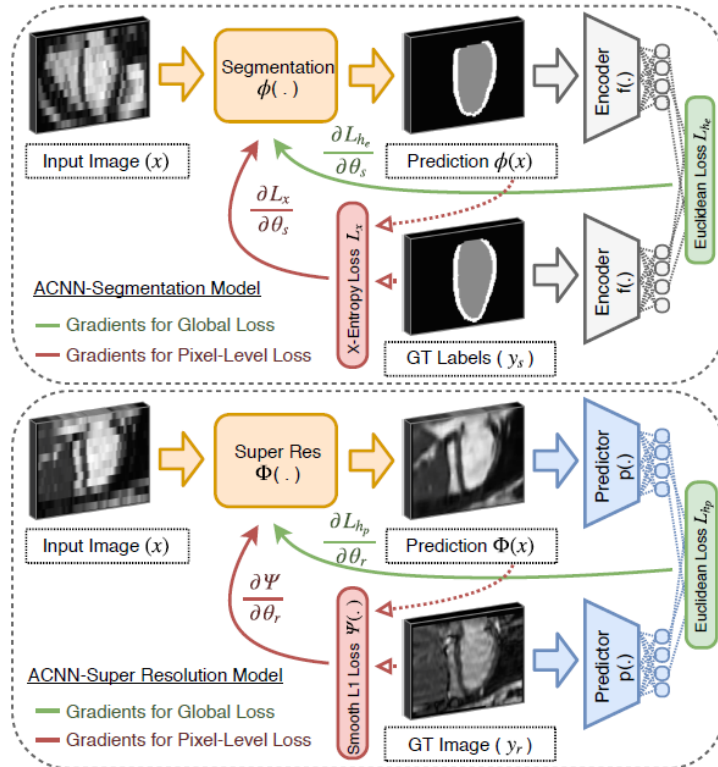


Figure 6.15: The proposed T-L network segmentation and super resolution tasks, this figure is collected from (Oktay et al., 2018).

Oktay *et al.* come up with T-L network in (Oktay et al., 2018), which is used for image segmentation or a super resolution problem. They apply the proposed novel structure for

cardiac image segmentation and enhancement. The schema of T-L network for segmentation or super resolution is presented in Fig.6.15. They introduce two losses to integrate shape prior and improve the accuracy of estimation at pixel level. The block "Segmentation" and "Super Res" in Fig.6.15 enhance the similarity between the estimation and the reference at pixel-level, the correspondent loss is "X-Entropy Loss L_x " for segmentation, "Smooth L_1 loss" for super resolution problem. The shape prior is integrated via a perceptual loss. The idea is to non-linearly transform the estimation and the reference into a low dimension space and to penalize the dissimilarity between them. In SRGAN, the perceptual loss is considered in the feature space described by VGG, since VGG is a deep network, this feature space is relatively huge compared to the space described by encoder. The segmentation results obtained with the proposed T-L model outperforms other methods in the literature.

The deep learning methods have a vast field of application in medical image processing tasks (Litjens et al., 2017), such as classification, detection, segmentation, registration and so forth. As with the development of the application of deep learning in the super resolution problem, researches in medical image processing have been influenced by deep learning methods, and now propose new architectures which permit to integrate priors to boost the performance of the network and facilitate the follow-up analysis and research.

Chapter 7

A proposed network for our super resolution problem: FSRCNN-Residual

In the previous chapters, we investigated variational approaches such as Total Variation based regularization (Toma et al., 2014b, Peyrin et al., 2015, Li et al., 2017a) and dictionary learning methods in (Wang et al., 2012) for super resolution. The numerical results showed that the dictionary learning based approach outperforms TV in our application. Yet, one difficulty of the dictionary learning methods in super resolution is how to properly couple high and low resolution dictionaries. Different mapping functions have been discussed in (Yang et al., 2010a, Wang et al., 2012, He et al., 2013). However, these functions are constrained to describe simple mapping relations, which are not as efficient as deep learning methods when the mapping relations are very complex.

In the last chapter, a review was given to illustrate approaches based on deep learning network for super resolution problems and the application in medical image processing. In this chapter, we propose a FSRCNN-Residual network to solve the joint super-resolution/segmentation problem for HR-pQCT images. Even though FSRCNN is not the best architecture for the super resolution problem, it has less parameters and provides acceptable results. The residual network refines the segmentation. U-net (Ronneberger et al., 2015), V-net (Milletari et al., 2016) and Volumetric ConvNets (Yu et al., 2017) are three prevailing approaches today. These networks are composed of a downsizing part which enables the representation of the image with a huge number of very abstract low spatial dimension features and an upsizing part which enables to retrieve the initial spatial dimension and to use the learnt features to achieve the segmentation. Differing from the typical segmentation methods, the residual network doesn't abstract image features as (Ronneberger et al., 2015), it learns the residual information between the label images and the estimations.

The proposed network is firstly applied to solve a super resolution problem with up-sampling factor $r = 2$ and the results are compared with formerly studied methods. Then a minor modification is introduced so that the network solves a super resolution problem with upscale factor $r = 3.42$.

This chapter is organized as follows: Sect.7.1 introduces the proposed architecture, Sect.7.2 presents numerical experiments for the super resolution problem with upsampling factor $r = 2$, Sect.7.3 solves a super resolution problem with upsampling factor $r = 3.42$ via the proposed network. We end this chapter with discussion and conclusion.

Before introducing the architecture, we first recall some notations representing different types of images. Let $\mathbf{g} \in \mathbb{R}^N$ denote the 3D HR-pQCT low resolution images, $\mathbf{f}^* \in \mathbb{R}^{N'}$ represents the 3D high resolution ground truth obtained from μ -CT, $\mathbf{f}_b^* \in \mathbb{R}^{N'}$ is the binary image of $\mathbf{f}^*$. $\mathbf{f}$ and $\mathbf{f}_b \in \mathbb{R}^{N'}$ are the gray level and binarized estimated super resolution images, respectively. $\mathbf{f}_s \in \mathbb{R}^{N'}$ is the output of the network, which is optimized to approximate $\mathbf{f}_b^*$ but in gray level. The isotropic upsampling factor is r , thus $N' = r^3 N$ for 3D images.

Figure 7.1: Illustration of the proposed FSRCNN-Residual network. The former sub-network FSRCNN(Dong et al., 2016) tackles the super resolution task, the sequential sub-network Residual(Kim et al., 2016a) serves for segmentation purpose. The symbols $\mathbf{f}$ and $\mathbf{f}_s$ denote super resolution images and output images, respectively. The blue cubes denote features, orange cubes represent filters. The numbers above the blue cubes are the number of filters ($d_1, s, d_2,$), d_2 is the number of filters in the Residual sub-network which is not depicted on this figure. The numbers below orange cubes are the length of 3D filter edges. For instance, when the length of filter edge is 5, the size of filter is $5 \times 5 \times 5$.

FSRCNN is an extension of SRCNN for the super resolution problem. In this work, it is represented as an 8-layer network, as shown in Fig.7.1. At the first layer, the number of filters is d_1 , and the size of each filter is $5 \times 5 \times 5$, the activation function after the convolution filters is Parametric Rectified Linear Unit (Prelu). $\text{Prelu}(x) = x$ if $x \geq 0$, otherwise $\text{Prelu}(x) = ax$, where a is adapted during the training process.

corresponding to the number of filters in the FSRCNN sub network, with $d_1 > s$. The $deconv(1,9)$ is a deconvolution layer without any activation function. It performs an upscale operation.

FSRCNN displays an hourglass shape. The first layer abstracts low resolution structural information. The second layer reduces the channel dimensionality of features. The four consecutive layers in the middle of the network serve as a non linear mapping function. It transforms the shrunk low resolution features into another space where the high resolution features are represented in low channel dimension. The sixth layer expands these low dimensionality features to a high dimensional level. And the last deconvolution layer can be compared to a decoding layer, which estimates the super resolution image $\mathbf{f}$.

7.1.2 The Residual sub-network

The residual network in Fig.7.1 is dedicated to the segmentation task. It includes 3 convolution layers and is represented as $conv(d_2,3) \rightarrow conv(d_2,3) \rightarrow conv(1,3)$, where d_2 is the number of filters. The sigmoid activation function is used for 2-label classification, defined as: $sigmoid(x) = 1/(1 + e^{-x})$. Following the work in (Kim et al., 2016a), the three convolution layers learn the residual information to refine the segmentation.

7.1.3 The loss functions

The network parameters are optimized regarding the total loss $\mathcal{L}$ which consists of two partial losses $\mathcal{L}_1$ and $\mathcal{L}_2$:

$$\begin{aligned}\mathcal{L}_1(\mathbf{f}^*, \mathbf{f}) &= \frac{1}{MN'} \sum_i^M \|\mathbf{f}_i^* - \mathbf{f}_i\|_2 \\ \mathcal{L}_2(\mathbf{f}_b^*, \mathbf{f}_s) &= - \frac{1}{MN'} \sum_i^M \sum_j^{N'} (\mathbf{f}_{b_i,j}^* \log(\mathbf{f}_{s_i,j}) + (1 - \mathbf{f}_{b_i,j}^*) \log(1 - \mathbf{f}_{s_i,j})) \\ \mathcal{L}(\mathbf{f}^*, \mathbf{f}_b^*, \mathbf{f}, \mathbf{f}_s) &= \mathcal{L}_1(\mathbf{f}^*, \mathbf{f}) + \mathcal{L}_2(\mathbf{f}_b^*, \mathbf{f}_s)\end{aligned}\tag{7.1}$$

In fact, $\mathbf{f}_i^*$ denotes high resolution patches, $\mathbf{f}_i$ denotes the output of FSRCNN network. By minimizing the L_2 distance between these patches, the parameters in the network will be updated. The denominator N' is used for taking the average on voxel, the denominator M is the size of batch. The parameters in the network are updated via gradient descent method. There are different ways to determine the gradient. In this work, we use stochastic gradient descent method. The idea is to randomly pick up M patches, the gradient is estimated by taking the average of these patches.

The explanations are the same for the denominator MN' in $\mathcal{L}_2$. $\mathcal{L}_2(\mathbf{f}_b^*, \mathbf{f}_s)$ corresponds to the cross entropy between the final output of the network and the binarized ground truth. For instance, when the voxel on the ground truth is 1, the function is reduced to $\mathbf{f}_b^* \log(\mathbf{f}_s)$, By maximizing this term, the $\mathbf{f}_s$ tends to be large value. When the voxel on the ground truth is 0, the function is reduced to $(1 - \mathbf{f}_b^*) \log(1 - \mathbf{f}_s)$. Maximizing this function pushes $\mathbf{f}_s$ to be small. The last but not the least, the activation function before the output of the network is a sigmoid function since the combination of cross entropy and sigmoid activation is a typical choice for 2-label classification tasks (LeCun et al., 2015).

7.2 Numerical experiments

7.2.1 Dataset and criteria

Our dataset is composed of 13 cadaveric samples including 7 radius and 6 tibias. The number of samples in the training/validation/test sets are 5/3/5, respectively. The low resolution images $\mathbf{g}$ are obtained from HR-pQCT with typical *in vivo* protocol, with $82\mu\text{m}$ isotropic voxels. The high resolution images are generated with $\mu\text{-CT}$ system, with voxel size $24\mu\text{m}$. After registration via Elatix (Klein et al., 2010, Shamonin et al., 2014), the voxel size of the ground truth $\mathbf{f}^*$ is $41\mu\text{m}$. $\mathbf{f}_b^*$ is obtained with Otsu segmentation of $\mathbf{f}^*$. The proposed method is applied to the selected VOI of $200 \times 200 \times 200$ voxels in the full images. The network is optimized via Adam (Kingma and Ba, 2014), different rotations are considered in data augmentation. The network is optimized with back propagation gradient descent method. At each iteration, the gradient is determined by averaging a number gradients obtained with different image patches. The number of these different image patches is named as batch size. In this work, the batch size is 16, $d_1 = 128$, $s = 24$, $d_2 = 64$. The code is implemented with Tensorflow, NVIDIA TITAN Xp, CUDA 9.0.

The next subsection presents super resolution results and compares the performance of different methods on small crops extracted from the trabecular samples.

7.2.2 Super resolution results for crops

Similarly to previous work, the Dice index, BV/TV and Conn.D have been considered to evaluate the reconstructed images.

Fig.7.2 illustrates gray level 2D slices of 3D trabecular micro-architecture restored with different methods. The comparison of gray level images shows that the proposed method outputs a quasi binary image with smooth and uniform structures. It's noteworthy that this crop is cut from the training set to show the best performance of the proposed approach.

The Fig.7.3 presents the corresponding binary images of Fig.7.2. The image given by TV regularization is smoother than low resolution, but the segmented structure is thicker than the high resolution image. CH+TV (Li et al., 2017a) is a variational approach solving nonconvex and nonsmooth optimization. It is a post-processing of TV, which is expected to provide structures as thin as the high resolution image. However, it fails to recover structural details. 2.5D SCDL is a dictionary learning approach based on (Wang et al., 2012), it manages to recover more details compared with CH+TV. The image estimated via the proposed method has a visual aspect comparable to the one of 2.5D SCDL, but with smoother borders.

2.5D SCDL is a dictionary learning based approach adapted from the work (Wang et al., 2012), it improves Dice and Conn.D except for BV/TV. It can be seen in Fig.7.2 that 2.5D SCDL recovers more details compared with CH+TV.

Both FSRCNN and the proposed method have improved BV/TV and Dice index, the relative error in term of Conn.D are comparable.

Tab.7.2 quantitatively compares different approaches on the training volumes whose 2D slices are shown on Fig.7.2. The large distance between the parameters measured from low resolution and high resolution explains the motivation of this super resolution/segmentation study. The TV regularization (Toma et al., 2014b) improves the BV/TV and

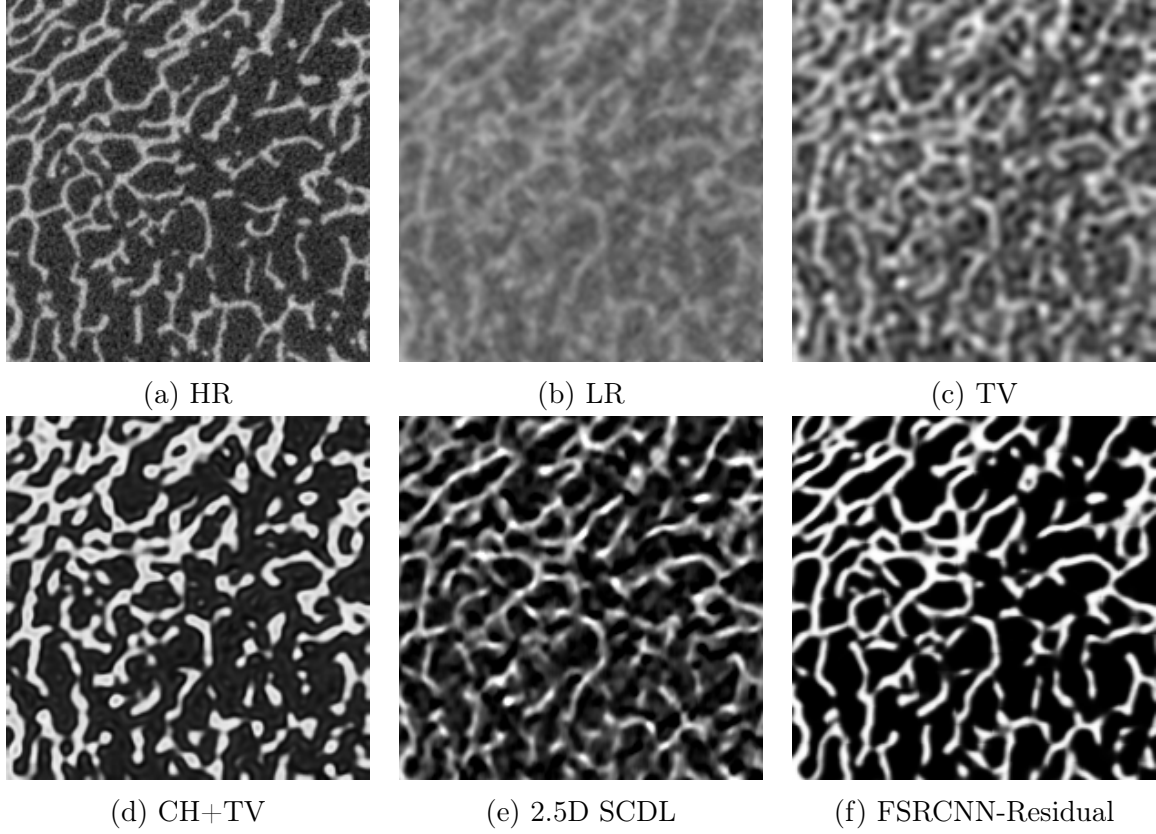


Figure 7.2: Illustration of a 2D slice of 3D image of the training set in gray level, the upsampling factor is 2. The label of the sample is 'C0008233-D0000812'. On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed FSRCNN-Residual, respectively.

Table 7.1: Comparison of different super resolution approaches. The volume is a $200 \times 200 \times 200$ crop cut from the sample 'C0008233-D0000812' in the training set, the upsampling factor is 2.

image	Dice	BV/TV(%)	Conn.D (mm^{-3})
LR		40.1	5.169
HR	1	25.6	3.713
TV	0.742	39.0	4.253
CH+TV	0.779	29.0	4.441
2.5D SCDL	0.797	29.7	3.821
FSRCNN	0.764	26.0	3.876
FSRCNN-Residual	0.808	29.1	3.546

Conn.D, but they are still far away from the parameters measured on HR. CH+TV (Li et al., 2017a) is a variational approach solving nonconvex and nonsmooth optimization. It has largely improved the Dice index and BV/TV, but the Conn.D is still limited and its corresponding image displayed in Fig.7.2 is not as well restored as expected: too many de-

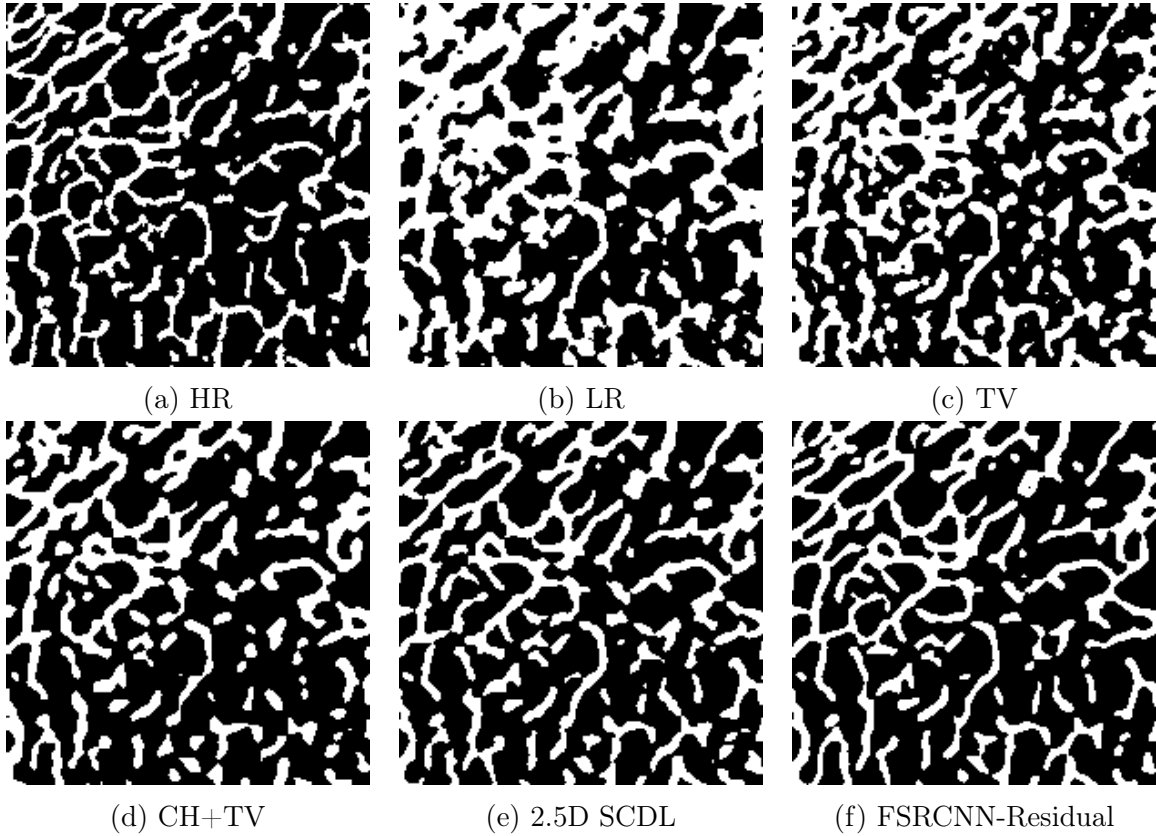


Figure 7.3: Illustration of a 2D slice of 3D image of the training set in binary, the up-sampling factor is 2. The label of the sample is "C0008233-D0000812". On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.

tails are lost. 2.5D SCDL is a dictionary learning based approach adapted from the work (Wang et al., 2012). Compared with CH+TV, 2.5D SCDL improves Dice and Conn.D but not BV/TV. It can be seen in Fig.7.2 that 2.5D SCDL recovers more details compared with CH+TV. FSRCNN gives the best BV/TV along with a comparable Conn.D to FSRCNN-Residual. The best Dice is obtained with the proposed FSRCNN-Residual network.

In order to see the generalization properties of the proposed method, different approaches have been applied on another crop cut from the test set. Fig.7.4 and Fig.7.5 present the gray level and binary 2D slices obtained with different methods. Tab.7.2 quantitatively compares different super resolution approaches on the test crop shown in Fig.7.5. Similar conclusion has been drawn: the proposed FSRCNN-Residual network visually outperforms all the previously investigated approaches. Although the Conn.D is not improved, the Dice index and BV/TV outperforms all the other approaches.

7.2.3 Discussion

Compared with FSRCNN-Residual network, 2.5D SCDL needs to learn three dictionaries for super resolution reconstruction. It takes a longer time to learn the dictionaries.

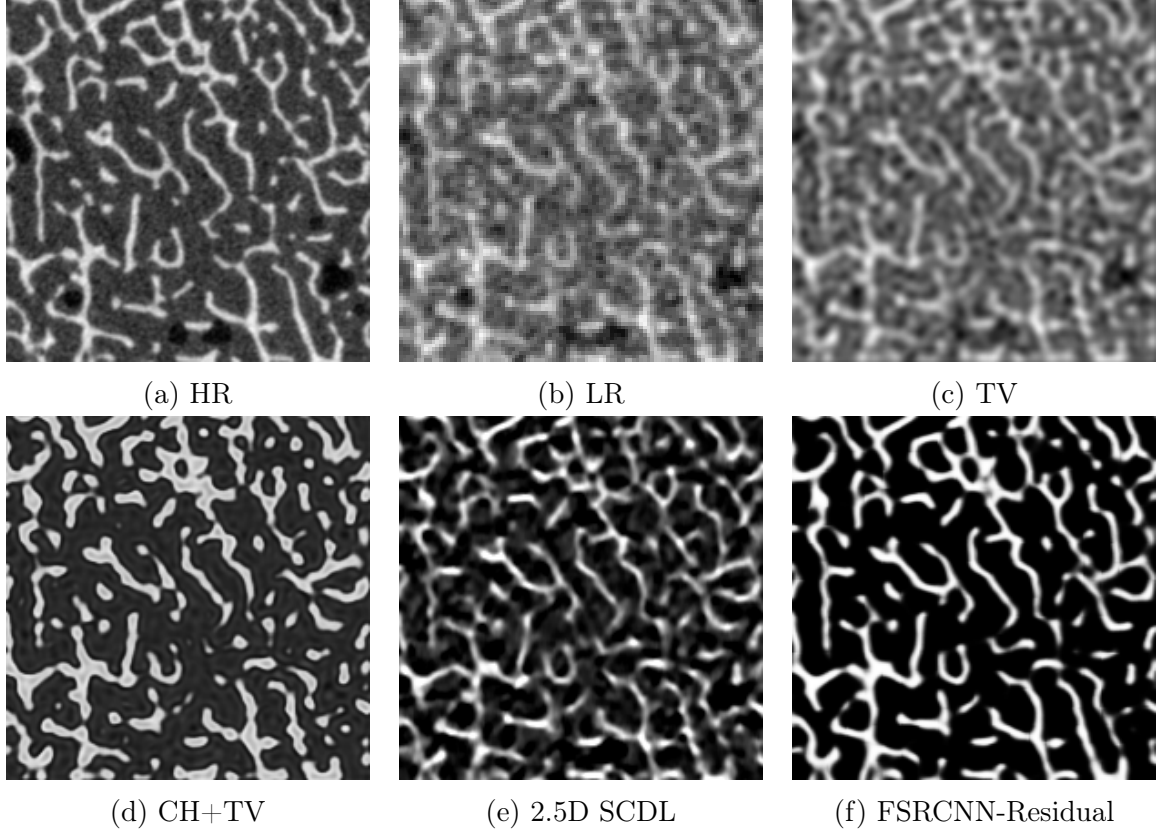


Figure 7.4: Illustration of a 2D slice of 3D image of the test set in gray level, the up-sampling factor is 2. The label of the sample is 'C0008231-F0000814'. On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.

Table 7.2: Comparison of different super resolution approaches. The tested volume is a $200 \times 200 \times 200$ crop cut from the sample 'F0000814-C0008231' in the test set.

image	Dice	BV/TV(%)	Conn.D (mm^{-3})
LR		36.0	4.496
HR	1	24.1	4.879
TV	0.690	35.2	3.903
CH+TV	0.684	25.2	4.103
2.5D SCDL	0.751	26.4	3.776
FSRCNN	0.822	24.4	3.932
FSRCNN-Residual	0.826	24.9	3.972

Furthermore, in terms of reconstruction time, FSRCNN-Residual spends less time than 2.5D SCDL. For these reasons, FSRCNN-Residual is a feasible choice to reconstruct super resolution images of entire HR-pQCT trabecular low resolution images.

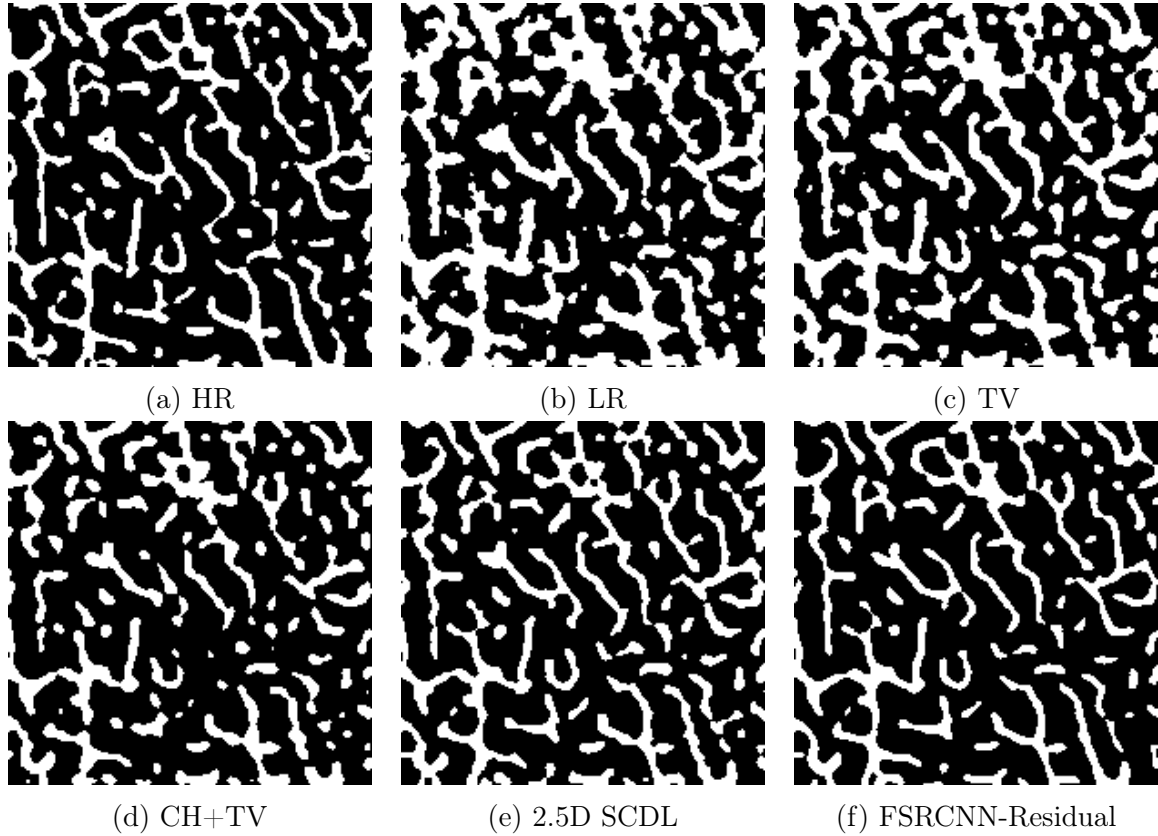


Figure 7.5: Illustration of a 2D slice of 3D image of the test set in binary, the upsampling factor is 2. The label of the sample is "C0008231-F0000814". On the first row, from left to right, high resolution, low resolution, TV super resolution image; on the second row, the images obtained with CH+TV, 2.5D SCDL and the proposed approach, respectively.

7.2.4 Conclusion

To sum up, we proposed a FSRCNN-Residual network to solve a joint super resolution/segmentation problem with upsampling factor $r = 2$. The two losses in the network inject priors of gray level and binary images so that the super resolution and segmentation tasks are performed at the same time. Our experimental results show that the proposed network outperforms previously studied approaches and is very promising. In the next section, we expect to increase the upsampling factor to 3.42 so that the voxel size of super resolution image turns to be $24\mu\text{m}$.

7.3 The application of FSRCNN-Residual in super resolution problem with upsampling factor 3.42

In this section, we attempt to improve the resolution of HR-pQCT images from the voxel size $82\mu\text{m}$ to $24\mu\text{m}$. In the literature, various networks have been proposed to solve super resolution problem with upsampling factor at 2, 3, or 4. In our dataset, the images with the best resolution is scanned from $\mu\text{-CT}$ and has a voxel size equal to $24\mu\text{m}$, the low resolution images are obtained from HR-pQCT with voxel size at $82\mu\text{m}$. From $82\mu\text{m}$ to

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

24 μm , the upsampling factor is around 3.42. Inasmuch that this factor is not integer, an easy way to improve the resolution of HR-pQCT images is to register the raw $\mu\text{-CT}$ images so that their voxel size turns to be 41 μm , therefore a super resolution problem with factor 2 is raised. All our previous work is based on this model.

However, the ground truth generated by registration is not as ideal as the primitive $\mu\text{-CT}$ images. Thanks to CNN, the convolution operation allows to output images with flexible size. Therefore, a super resolution problem from voxel size 82 μm to 24 μm is much simpler.

In this section, we first introduce the adapted architecture, then present implementation details and show reconstructed super resolution crops. Afterward, a morphological analysis over entire samples will be given.

7.3.1 The adapted FSRCNN-Residual

If we consider image pairs with low resolution images at 82 μm voxel size and high resolution images at 24 μm , 12 low resolution voxels have the same length as 41 voxels at high resolution scale, for $12 \times 82\mu\text{m} = 41 \times 24\mu\text{m}$.

With these image pairs, the low resolution patches are fed into the network as input images of size $12 \times 12 \times 12$, the correspondent high resolution patches of size $41 \times 41 \times 41$ (in gray level and binary) are fed into the network as references f^* and f_b^* .

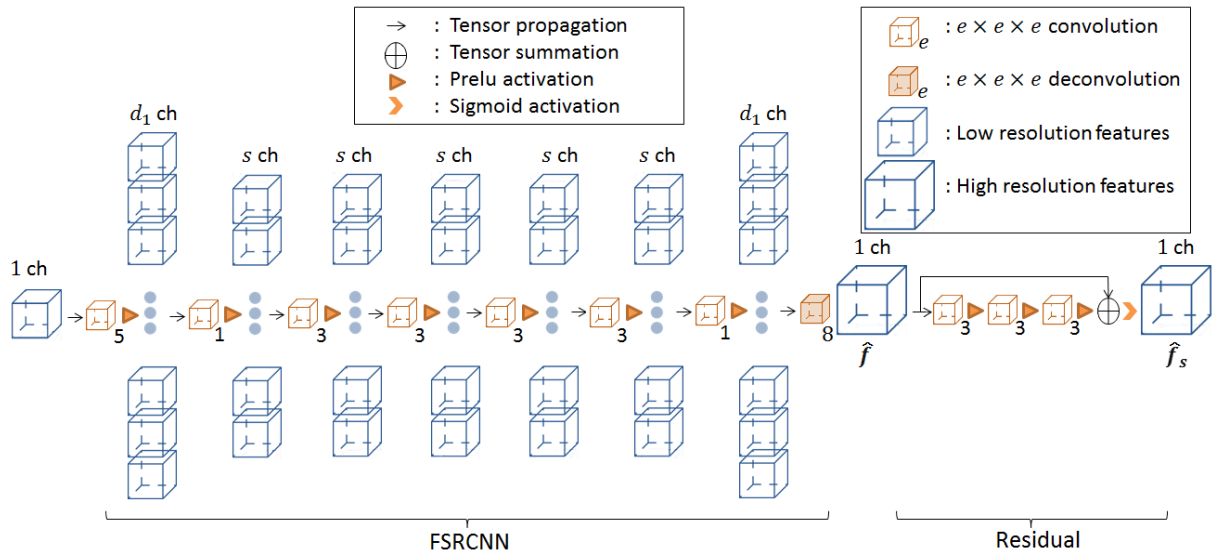


Figure 7.6: Illustration of the proposed FSRCNN-Residual network. The former sub-network FSRCNN(Dong et al., 2016) tackles the super resolution task, the sequential sub-network Residual(Kim et al., 2016a) serves for segmentation purpose. The symbols f and f_s denote super resolution images and output images, respectively. The blue cubes denote features, orange cubes represent filters. The numbers above the blue cubes are the number of filters (d_1, s, d_2), d_2 is the number of filters in the Residual sub-network which is not depicted on this figure. The numbers below orange cubes are the length of 3D filter edges. For instance, when the length of filter edge is 5, the size of filter is $5 \times 5 \times 5$. It's noteworthy that the size of the deconvolution filter is $8 \times 8 \times 8$, stride of convolution is 3, no padding involved.

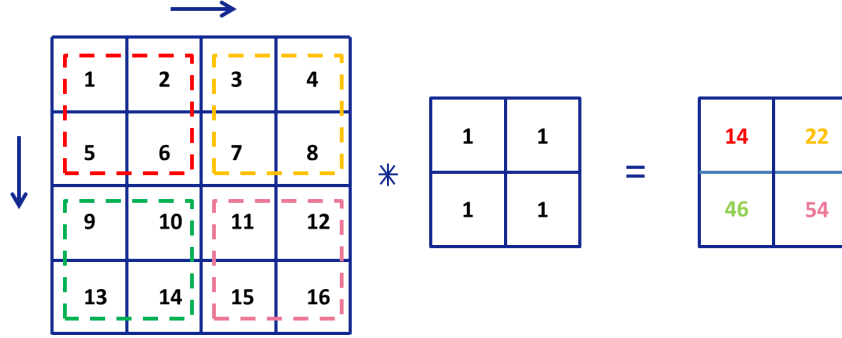


Figure 7.7: Demonstration of convolution with stride: a 4×4 signal is convolved with a 2×2 filter with stride 2.

Fig.7.6 draws the architecture of the network solving the super resolution problem with factor 3.42. All the first 7 layers in FSRCNN keep the same form as Sect. 7.1. But the deconvolution layer needs to be changed.

In 1D signal convolution, if the length of signal is L_{in} , the filter size is e , padding factor is p , stride is s , the output length L_{out} is determined by

$$L_{out} = \lfloor \frac{L_{in} - e + 2p}{s} \rfloor + 1 \quad (7.2)$$

where $\lfloor \cdot \rfloor$ takes the max integer which is not larger than the input. In the following, we give an example to explain the convolution with stride.

If 1D signal is of length 4, convolution with a 1D signal of length 2, with stride 2, no padding is considered, then we have

$$L_{out} = \lfloor \frac{4 - 2 + 2 * 0}{2} \rfloor + 1 = 2 \quad (7.3)$$

Similarly, in our problem, if a 1D signal at length 41, is convolved with a filter at length 8, the stride is 3, no padding is considered, then we obtain:

$$L_{out} = \lfloor \frac{41 - 8 + 2 * 0}{3} \rfloor + 1 = 12 \quad (7.4)$$

This is a forward convolution direction, which can be interpreted as a down sampling process. If we perform the inverse operation, the operator turns to be an upsampling operation, where the upsampling factor is 3.42.

To sum up, to solve a super resolution problem with upsampling factor 3.42, we set the size of deconvolution filter to 8, stride is 3, padding factor is 0.

7.3.2 Numerical results

Reconstructed super resolution crop

Fig.7.8 presents 2D slices of 3D crops at voxel size $24\mu m$. The crop is cut from a trabecular sample in the test set. It can be seen that the structures on the resolved super resolution image are very similar to the ground truth. More details have been recovered

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

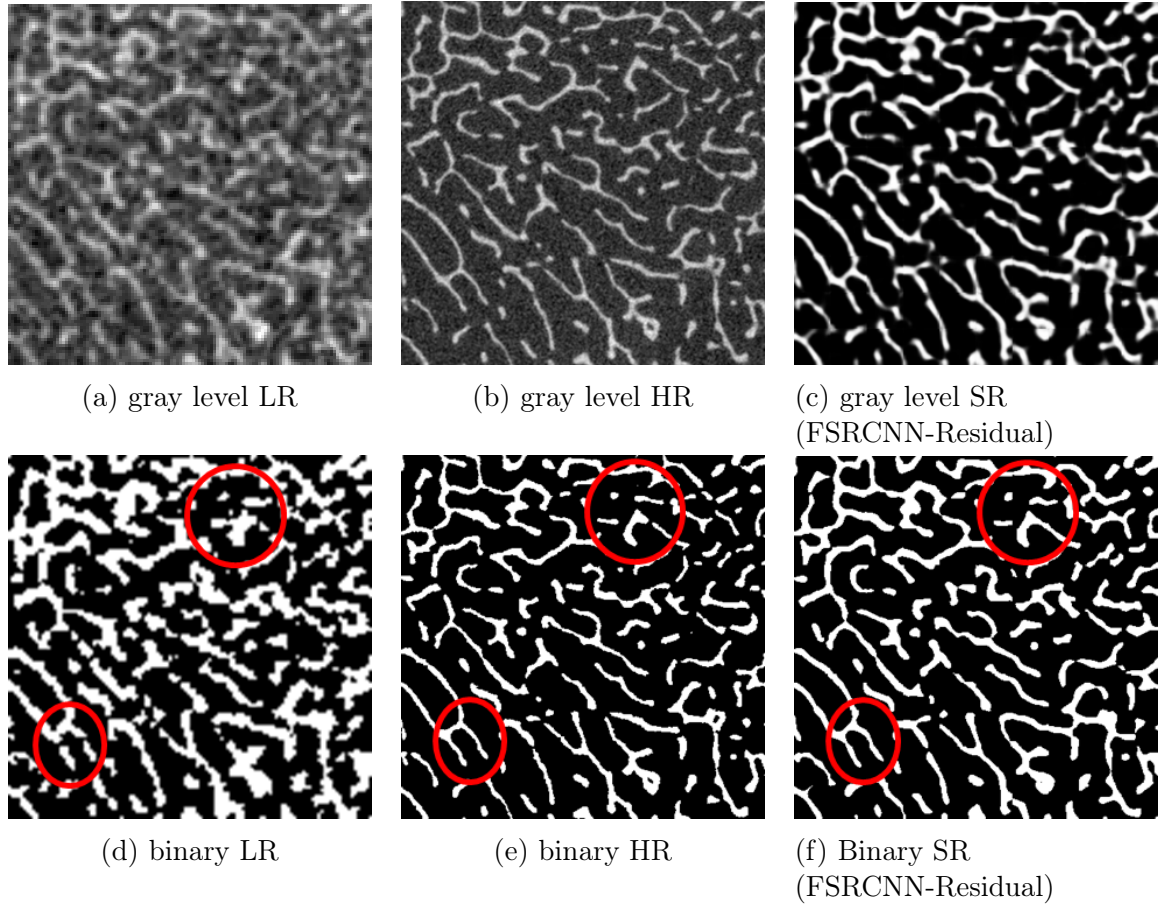


Figure 7.8: Illustration of a 2D slice of 3D volume at voxel size $24\mu\text{m}$. The crop is cut from test set. The label of sample is "C0008230-G0000818". The images on the top row are gray level images, from left to right: low resolution, high resolution, our method; All the images on the bottom row are binary images segmented with Otsu, from left to right: low resolution, high resolution and our proposed method.

Table 7.3: Comparison of different super resolution approaches. The evaluated subject is a $410 \times 410 \times 123$ crop cut from the sample 'C0008230-G0000818' in the test set.

image	Dice	BV/TV(%)	Conn.D (mm^{-3})
LR		32.9	4.839
HR	1	18.9	4.230
SR	0.817	23.2	3.086

from the low resolution images. For instance, the structures within red circles highlight the improvement brought by our proposed approach.

Tab.7.3 quantifies the differences of the volumes whose slices are presented in Fig.7.8. The BV/TV has been improved, and the index Dice indicates that the obtained super resolution image is an acceptable approximation to the ground truth. Nevertheless, the proposed method failed to preserve the Conn.D of the low resolution images.

Dataset

The dataset is composed of 13 cadaveric samples including 7 radius and 6 tibias. The number of samples in the training/test sets are 4/9, respectively. At current stage, we did not refine the parameters. The 4 samples in the training set include 2 tibia and 2 radius samples. The voxel size of high resolution images in the dataset is $24\mu\text{m}$, binary high resolution images are generated with Otsu segmentation of high resolution images. The low resolution images $\mathbf{g}$ are obtained with HR-pQCT with in-vivo protocol, the voxel size is $82\mu\text{m}$.

Reconstructed super resolution samples

We perform this super resolution method on entire samples to confirm the conclusion and perform morphological analysis. Fig.7.9-7.11 illustrates one slice of the entire samples at low resolution, high resolution and after super resolution.

The sample in Fig.7.9 is a tibia sample in the training set. It represents the best performance of the network. The trabecular micro architecture of the binary super resolution image is as delicate as that of the high resolution image. Even for the cortical bone regions, the structures are very similar.

The sample in Fig.7.10 is a tibia sample picked up from the test set. The super resolution image on this figure shows that the trained network keeps its efficiency on test set images. The trabecular micro architecture of the super resolution image has been recovered from low resolution images and the thickness of trabeculea is visually comparable to that of high resolution images. Moreover, the cortical regions have been ameliorated as well. Yet, the cortical regions on super resolution images are not as porous as that on high resolution images. It can be explained by the fact that the quantity of patches containing cortical bones is small in the training set, meanwhile variations on the cortical regions may fade out because of the normalization over the entire images.

Fig.7.11 reveals a flaw of the proposed network that we can't ignore. When the intensity on the bone image is not uniform, the normalization during the data preprocessing may improve the extraction of the structure information on the regions with weak contrast. As a result, the super resolution image loses bone mass and some structural details disappear.

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

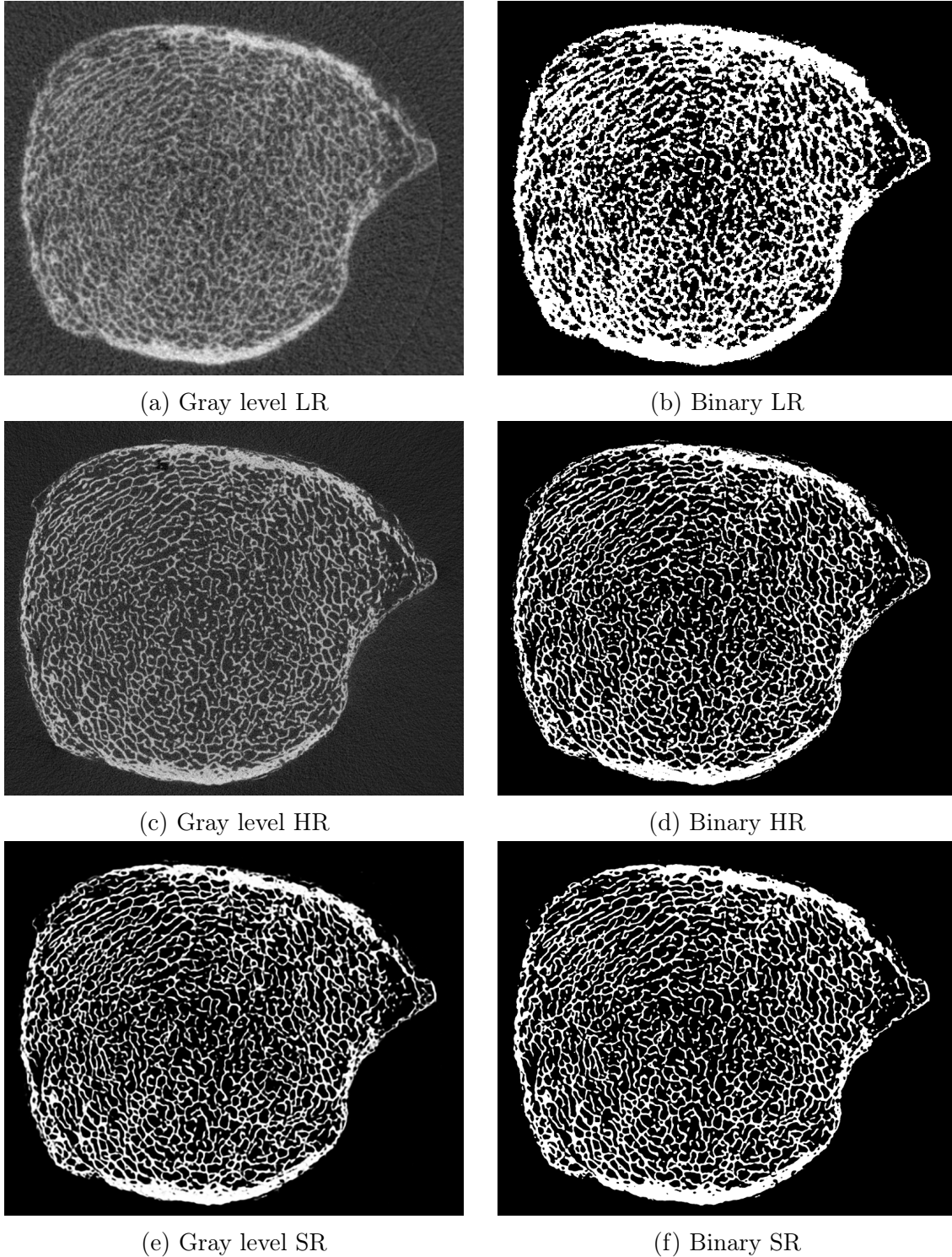


Figure 7.9: Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The sample is a tibia in the training set. The label of the sample is "C0008233-D0000812". This figure is used to show the best performance of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.

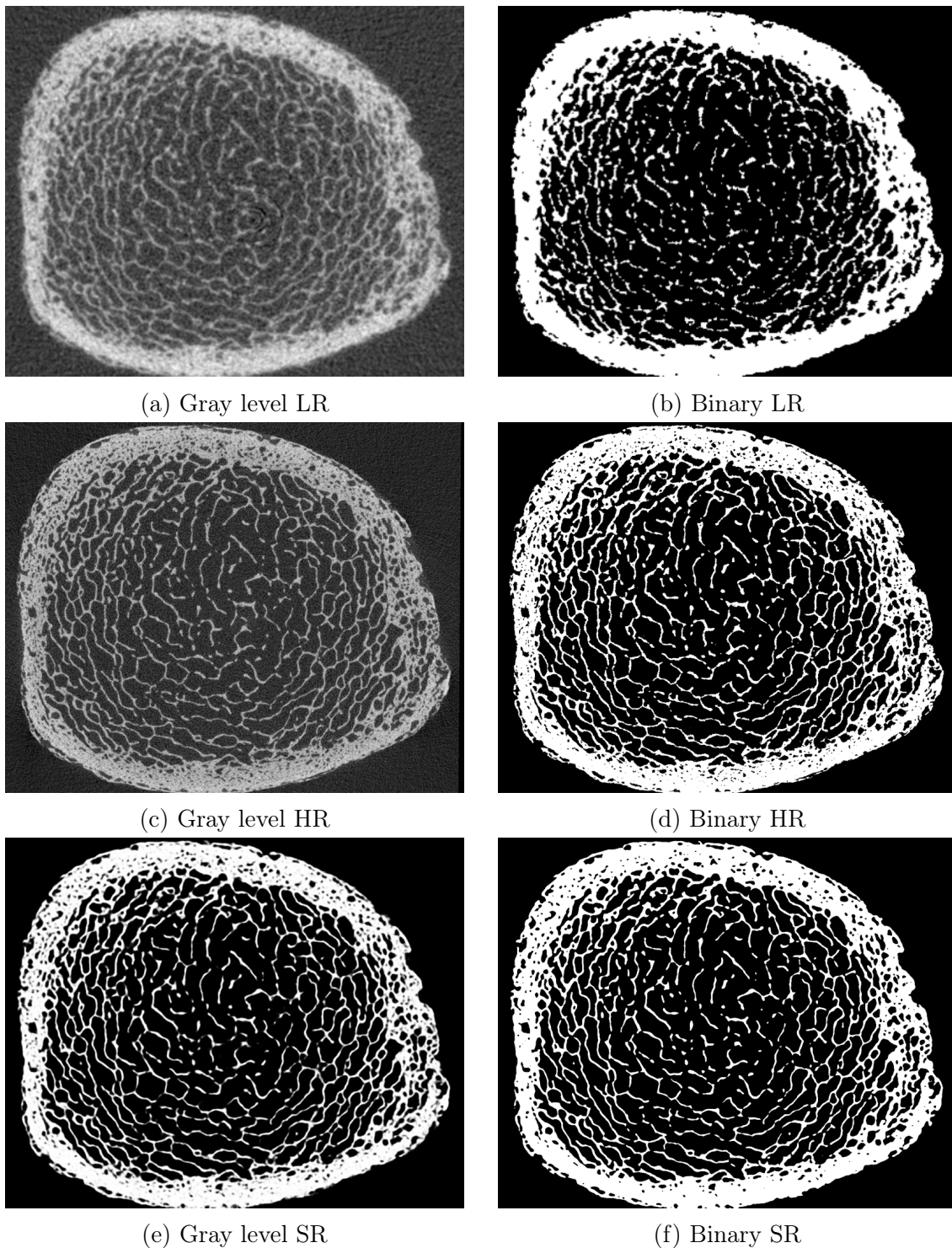


Figure 7.10: Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The sample is a tibia in the test set. The label of the sample 'C0008233-G0000812'. This figure is used to show the generalization property of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

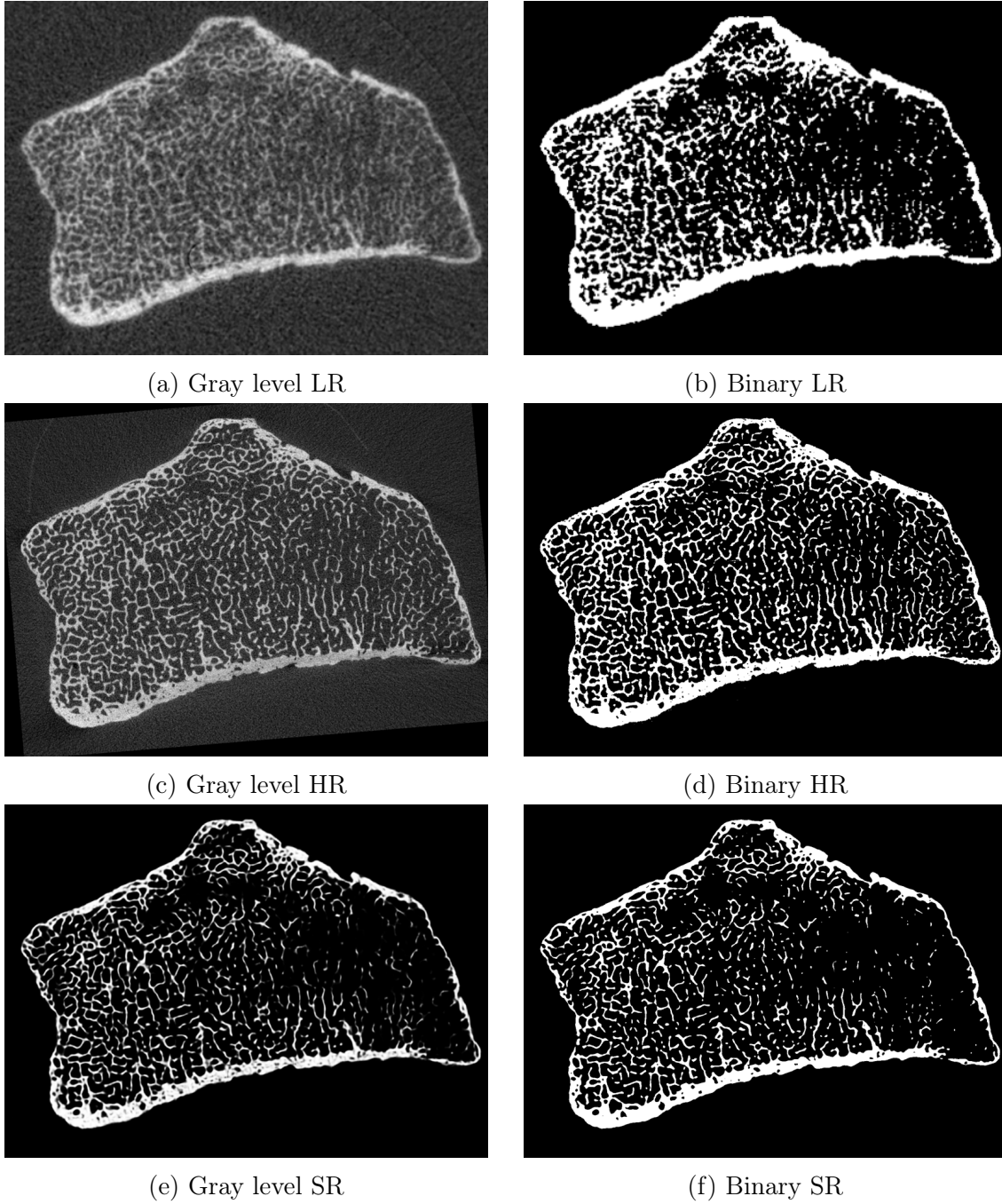


Figure 7.11: Illustration of a 2D slice of 3D entire sample at voxel size $24\mu\text{m}$. The tested sample is a radius in the test set. The label of the image is 'C0008231-F0000814'. This figure is used to show the generalization property of the network. The images on the left column are gray level images, from top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method; all the images on the right column are binary images. They are segmented with Otsu threshold, from top to bottom: low resolution, high resolution and super resolution image reconstructed via our method.

7.3.3 Computation trabecular bone structure parameters

The aim of this section is to compute quantitative trabecular bone structure parameters on the whole 3D images, before and after applying the super resolution processing. The images in the dataset contain trabecular and cortical bone regions. Nevertheless, the cortical bone is not necessary for trabecular bone analysis. In this section, we explain how to generate the mask to remove cortical region.

Briefly speaking, a mask is first generated for a high resolution sample. Afterward, the mask for the correspondent low resolution sample is obtained based on the mask at high resolution scale. In the following, we're going to detail how to generate both masks.

Mask of high resolution images

First of all, the high resolution images are sequentially eroded and dilated by mathematical morphology so that a rough cortical region is estimated. It's noteworthy that the number of erosion and dilation varies with cortical bone width.

Then, at each slice, we connect the two farthest points horizontally and vertically to generate a primitive mask. When the border of the sample is not convex, it is necessary to split the 2D slice into sub blocks so that the masks at sub blocks are convex. This step ensures that only the bone regions are considered.

Finally, the masks are shrunk by mathematical morphology so that it only covers the trabecular region.

Mask of low resolution images

The mask of low resolution images are generated via convolution operation. In the first step, the mask of high resolution images is cut into a grid, every patch has a size $41 \times 41 \times 41$. Then, each cell is convolved with an average filter of size $8 \times 8 \times 8$, with stride 3. Therefore, each high resolution patch outputs a small patch of size $12 \times 12 \times 12$ at low resolution scale. The final masks of low resolution images are obtained via Otsu segmentation of the images which are spliced by the patches at low resolution scale.

Fig.7.12 and Fig.7.13 illustrate the trabecular bone region extracted with the masks. The sample on Fig.7.12 is a tibia, on Fig.7.13 is a radius.

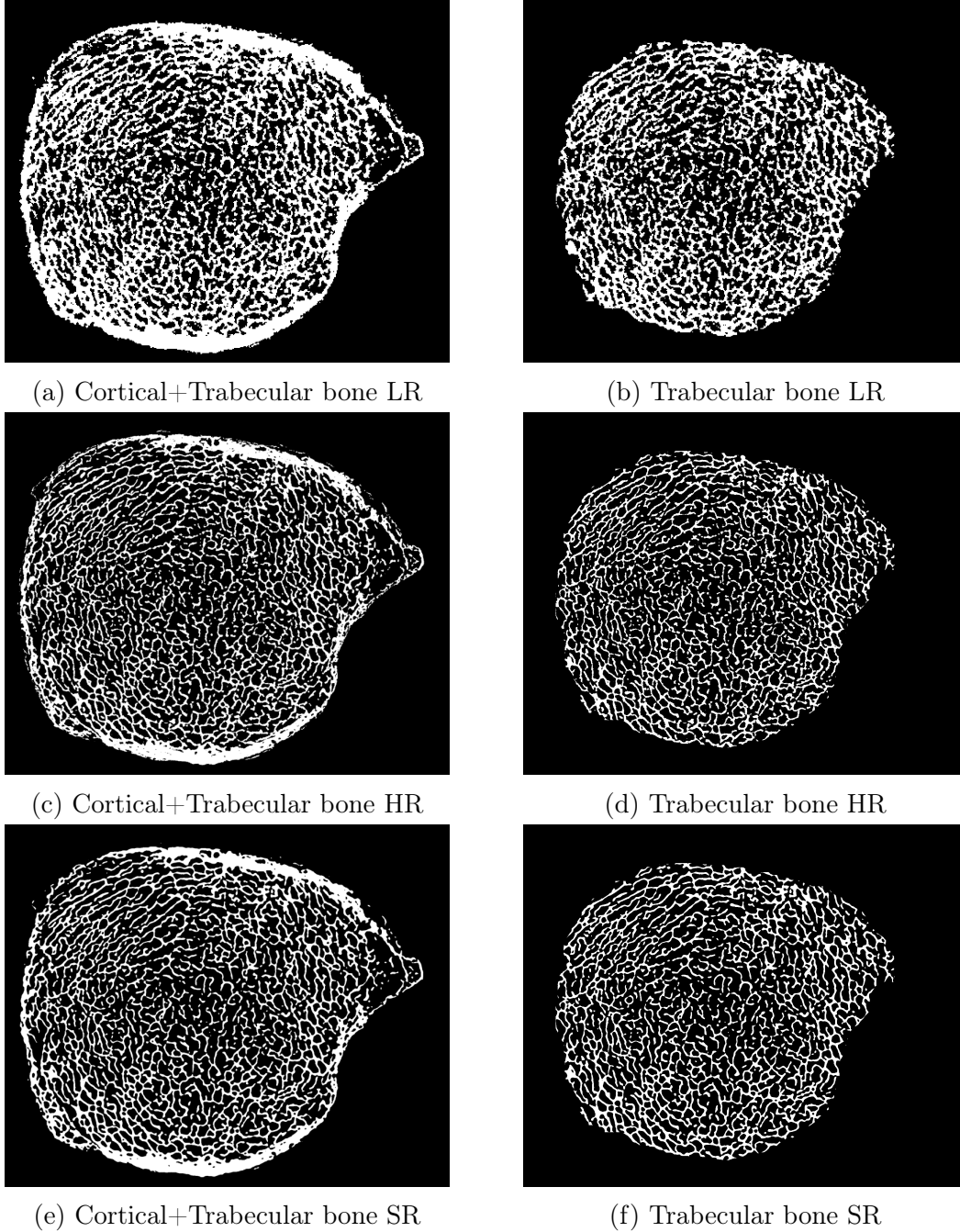


Figure 7.12: Illustration of the effects of mask on a tibia sample. The label of the sample is "C0008233-D0000812". The images on the left column include cortical as well as trabecular region. The images on right column only displays the trabecular structures abstracted by the mask. From top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method.

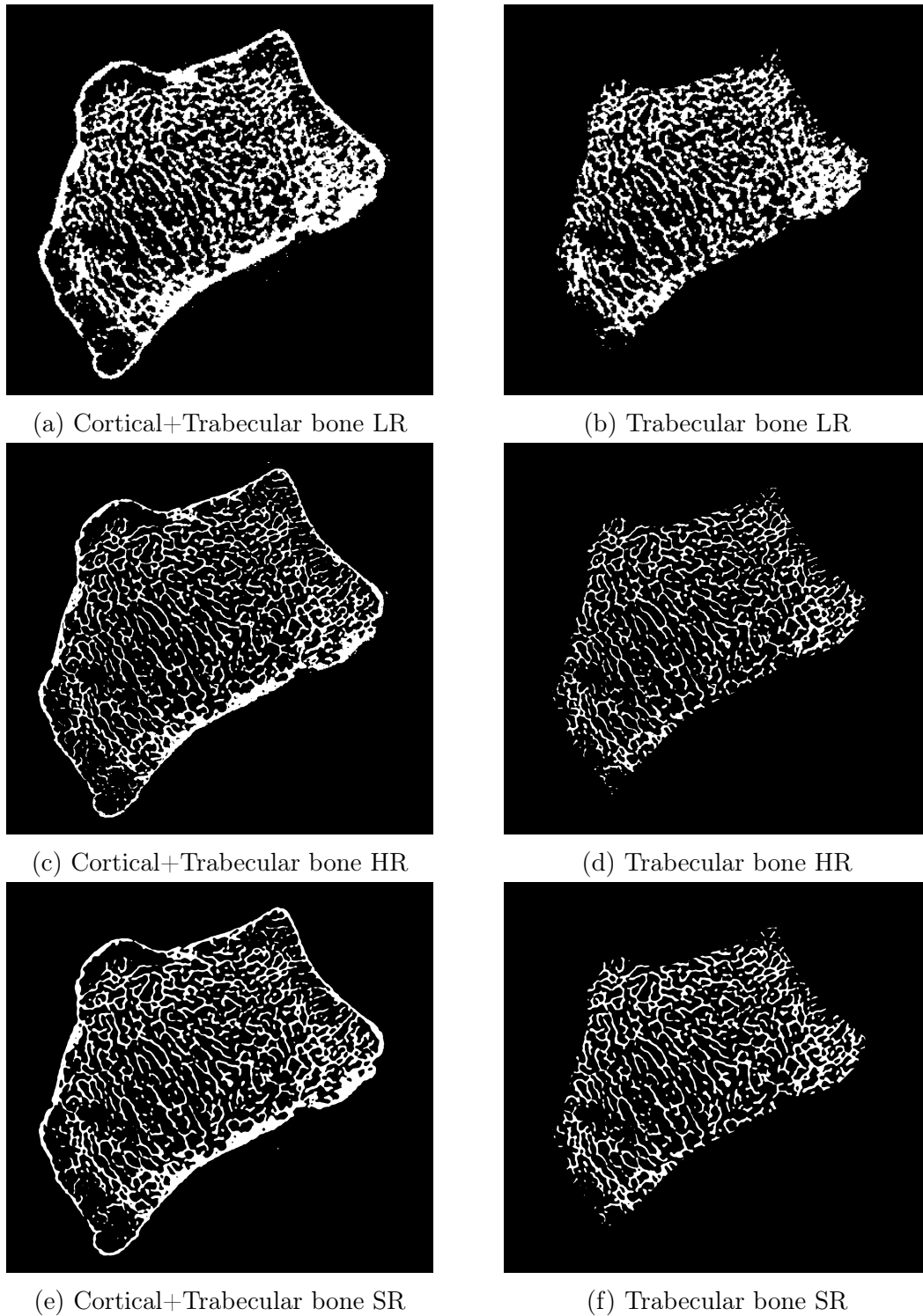


Figure 7.13: Illustration of the effects of mask on a radius sample. The label of the sample is "C0008230-G0000818". The images on the left column include cortical and trabecular region. The images on the right column displays the trabecular structures abstracted by the mask. From top to bottom: low resolution, high resolution, super resolution image reconstructed via our proposed method.

7.3.4 Image quantification

Since the volumes of 3D entire samples are very huge, here we summarize the BV/TV and the Dice index on the reconstructed samples in Tab 7.4. The Conn.D and other morphological analysis will be presented in the future. Among the 13 huge samples in the dataset, 9 of them have been successfully reconstructed with a good BV/TV and Dice. However, the other four samples have poor BV/TV estimation.

Fig. 7.15 presents 2D slices of sample 'C0008230-D0000818' and 'C0008233-F0000812'. As can be seen in the figure on the left column, some missing details on the low resolution images are recovered by the proposed network. Nevertheless, the bone structures are still thicker than the ground truth. Therefore, the BV/TV tends to be overestimated.

Meanwhile, the estimation of BV/TV on the samples 'C0008233-F0000812', 'C0008231-G0000814' as well as 'C0008231-F0000814' are underestimated, as shown on the right column of Fig. 7.15 and the two columns on Fig. 7.16. The super resolved images tend to be sparse, at the same time, some bone masses are visually shrinking or missing.

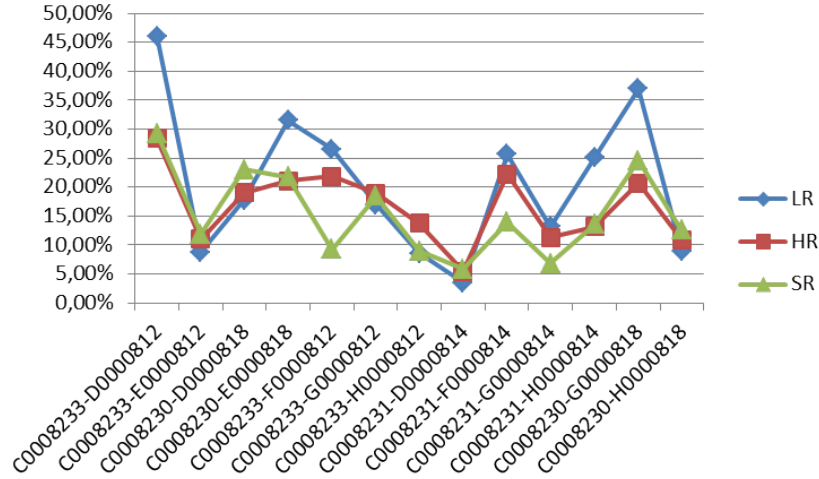


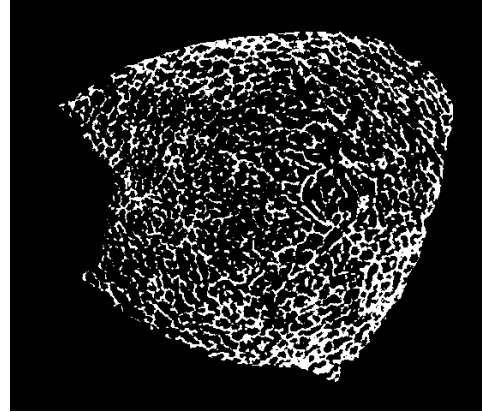
Figure 7.14: Summary of BV/TV over the entire dataset. The first 4 samples are in the training set.

Table 7.4: Summary of Dice and BV/TV over the entire dataset. The first 4 samples are in the training set, the other 10 are in the test set. The cells in yellow highlights the poor estimation of BV/TV on super resolution samples.

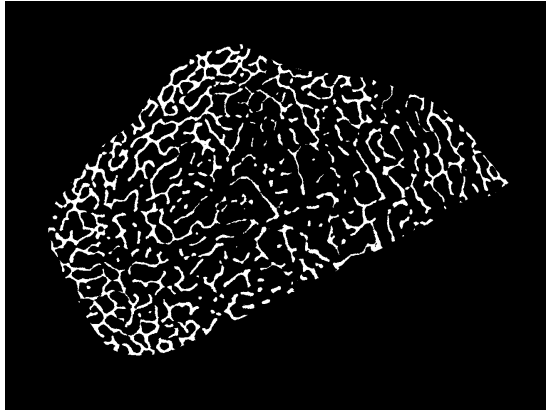
samples	DICE SR	BV/TV SR	BV/TV LR	BV/TV HR	BV/TV SR error	BV/TV LR error
C0008233-D0000812	0.846	29.29%	46.08%	28.37%	3.26%	62.45%
C0008233-E0000812	0.805	11.87%	8.80%	11.00%	7.91%	20.01%
C0008230-D0000818	0.831	23.05%	17.78%	19.04%	21.06%	6.59%
C0008230-E0000818	0.832	21.79%	31.47%	21.06%	3.45%	49.41%
C0008233-F0000812	0.558	9.22%	26.52%	21.88%	57.87%	21.18%
C0008233-G0000812	0.823	18.38%	16.93%	18.92%	2.90%	10.52%
C0008233-H0000812	0.710	8.94%	8.55%	13.70%	34.75%	37.59%
C0008231-D0000814	0.771	5.81%	3.46%	5.46%	6.36%	36.67%
C0008231-F0000814	0.703	13.89%	25.61%	22.28%	37.66%	14.98%
C0008231-G0000814	0.675	6.82%	13.17%	11.30%	39.64%	16.62%
C0008231-H0000814	0.763	13.53%	25.06%	13.10%	3.30%	91.29%
C0008230-G0000818	0.815	24.53%	37.02%	20.72%	18.42%	78.69%
C0008230-H0000818	0.776	12.60%	8.92%	10.94%	15.13%	18.47%



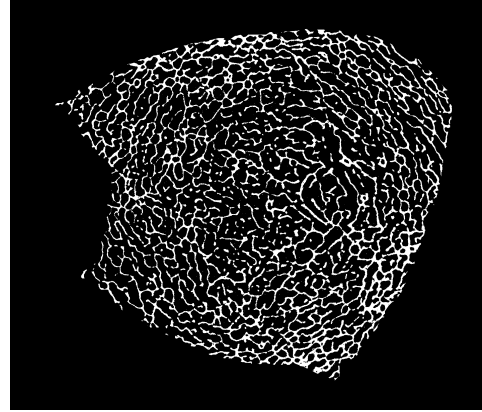
(a) 'C0008230-D0000818' Trabeculae LR



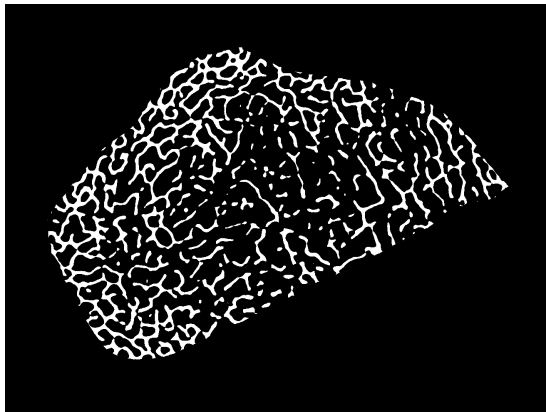
(b) 'C0008233-F0000812' Trabeculae LR



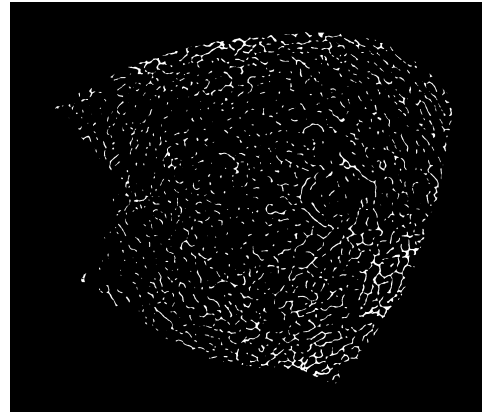
(c) 'C0008230-D0000818' Trabeculae HR



(d) 'C0008233-F0000812' Trabeculae HR

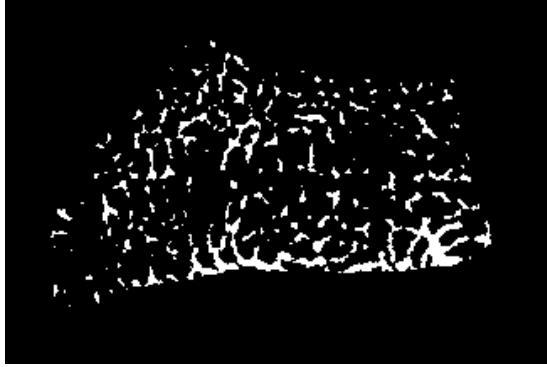


(e) 'C0008230-D0000814' Trabeculae SR

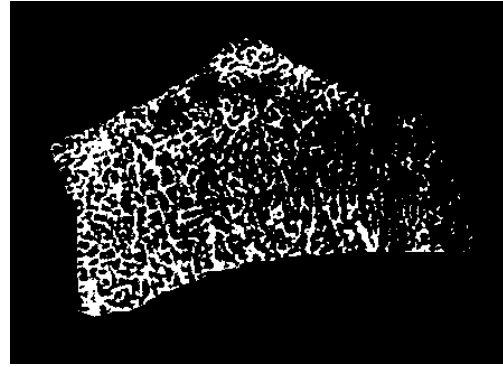


(f) 'C0008233-F0000812' Trabeculae SR

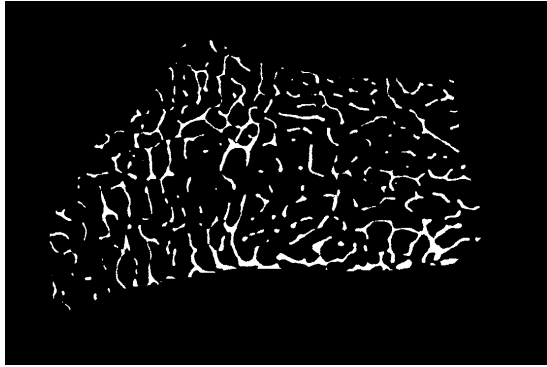
Figure 7.15: Illustration of samples with inferior BV/TV estimation. The trabecular samples on the left column is 'C0008230-D0000818', on the right is 'C0008233-F0000812'. From up to bottom, the images are binary low resolution, high resolution and super resolution images.



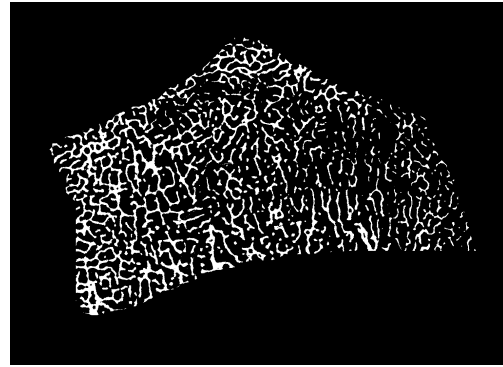
(a) 'C0008231-G0000814' Trabecular bone LR



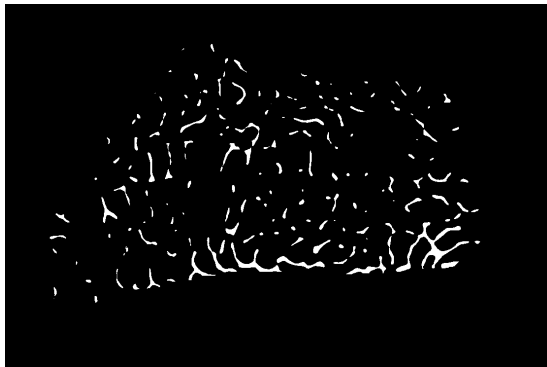
(b) 'C0008231-F0000814' Trabecular bone LR



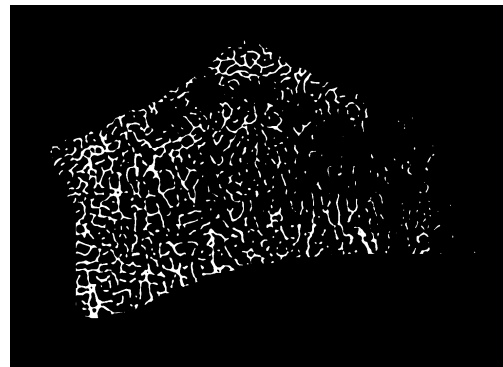
(c) 'C0008231-G0000814' Trabecular bone HR



(d) 'C0008231-F0000814' Trabecular bone HR



(e) 'C0008231-G0000814' Trabecular bone SR



(f) 'C0008231-F0000814' Trabecular bone SR

Figure 7.16: Illustration of samples with inferior BV/TV estimation. The trabecular samples on the left column is 'C0008231-G0000814', on the right is 'C0008231-F0000814'. From up to bottom, the images are binary low resolution, high resolution and super resolution images.

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

Overall, the performance of the proposed network varies with bone samples. 9/13 samples are well reconstructed in terms of BV/TV, however, 4/13 samples are not restored as well as expected. Generally speaking, there is a generalization problem of the network. More precisely, the statistic distributions of the samples are not homogeneous. In this primitive work, before being fed into the network, the images are normalized via simply dividing by the maximum value of the images. Yet, some high intensity noise scales the intensity of some bone structures on the dark region into even lower value, as a result, these structural information may be regarded as background by the network.

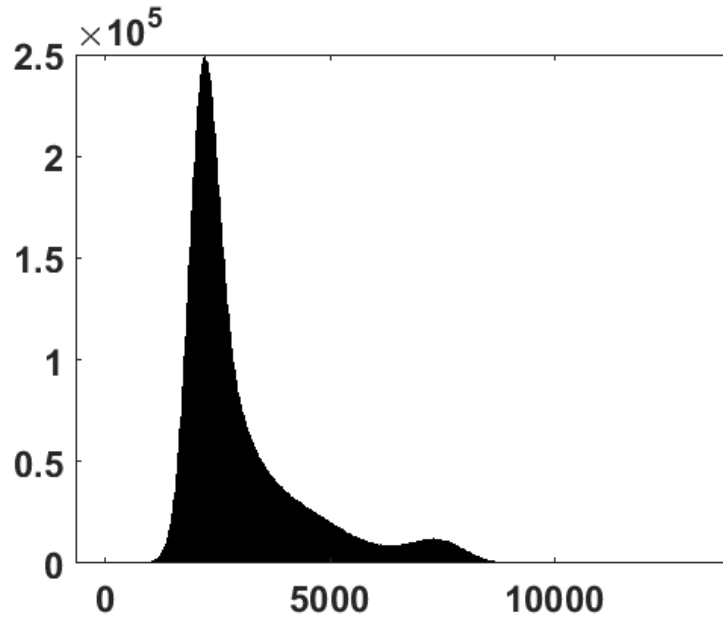


Figure 7.17: The histogram distribution of the sample 'C0008231-F0000814' at low resolution scale. It's noteworthy that this histogram display a bimodal shape since the cortical region is taken account of. There will be only one peak observed in the histogram if no cortical region is considered.

To prove this explanation, different scale factors have been applied for the normalization of the sample 'C0008233-F0000812'. The largest value on the sample 'C0008233-F0000812' at low resolution scale is 15773, its histogram is displayed on Fig.7.17. Besides 15773, 5000, 10000 along with 20000 are used to scale the low resolution image. Their correspondent super resolution images are illustrated in Fig.7.18, the Dice and BV/TV are summarized in Tab.7.5. Both Fig.7.18 and Tab.7.5 reveal that a larger scale factor leads to thinner or sparser results, a smaller scale factor favors the preservation of bone structure, when it is extremely small, even the background can be recovered as bone structure.

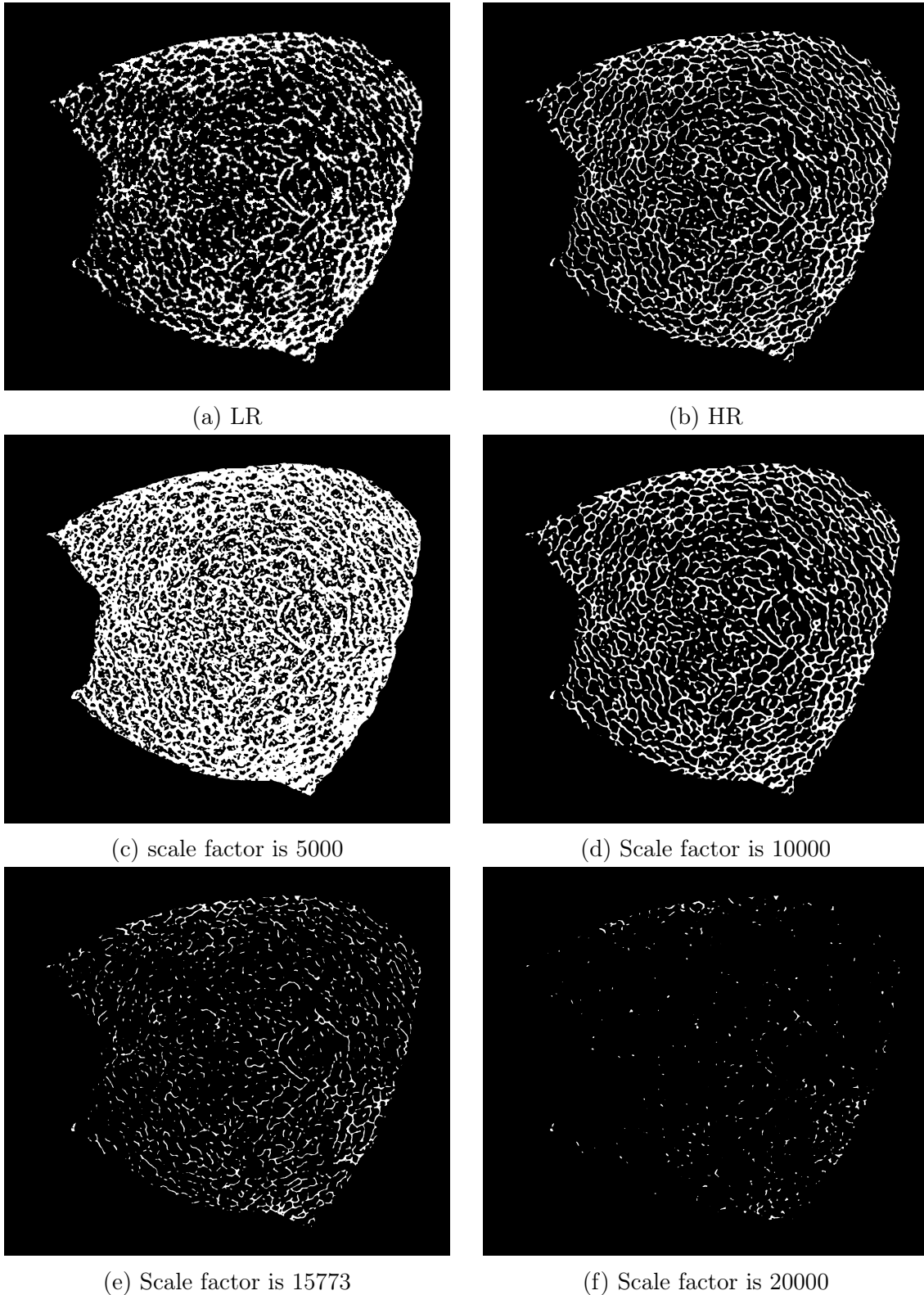


Figure 7.18: Comparison of super resolution images obtained with different scale factor in 'Normalization'. All the images are binary.

7.3. THE APPLICATION OF FSRCNN-RESIDUAL IN SUPER RESOLUTION PROBLEM WITH UPSAMPLING FACTOR 3.42

Table 7.5: Comparison of super resolution images obtained with different scale factors in the 'normalization'. The investigated sample is 'C0008233-F0000812'.

images	DICE	BV/TV
LR		26.52%
HR	1	21.88%
SR scale factor=5000	0.491	67.10%
SR scale factor=10000	0.819	24.49%
SR scale factor=15773	0.558	9.22%
SR scale factor=20000	0.233	3.06%

Table 7.6: Analysis of super resolution BV/TV error with the histogram of low resolution input images. The 'peak 1' and 'peak 2' are the two most obvious peaks positions (intensity) in the histogram of low resolution images. When the cortical regions are thin or non-homogeneous, the histograms do not have bimodal shape, in this case, the correspondent cells in the 'peak 2' column are marked with 'N', similarly for the column 'max/peak 2'. The column 'max' records the max intensity value on the low resolution image. 'max/-peak 2' reflects the relative position between the second peak in the histogram and the max value in the low resolution images. 'SR BV/TV error' is the relative error of the BV/TV of super resolution images with respect to the ground truth. When its sign is '+', it means that the BV/TV of the super resolution image is overestimated, the sign '-' means that the BV/TV is underestimated. Cells marked in orange highlight relevant information of samples whose super resolution BV/TV are highly underestimated.

Samples	peak 1	peak 2	max	max/peak 2	SR BV/TV error
C0008233-D0000812	2220	N	10795	N	3.26%
C0008233-E0000812	2255	7058	9926	1.41	7.91%
C0008233-F0000812	2319	7322	15773	2.15	-57.87%
C0008233-G0000812	2281	7180	11210	1.56	-2.90%
C0008233-H0000812	2242	7536	13539	1.80	-34.75%
C0008231-D0000814	2199	7260	10079	1.39	6.36%
C0008231-F0000814	2227	7256	13288	1.83	-37.66%
C0008231-G0000814	2280	N	12624	N	-39.64%
C0008231-H0000814	2266	N	9731	N	3.30%
C0008230-D0000818	2275	7546	9605	1.27	21.06%
C0008230-E0000818	2197	N	10573	N	3.45%
C0008230-G0000818	2153	N	9443	N	18.42%
C0008230-H0000818	2270	7088	9824	1.39	15.13%

Then we extend this finding to all the samples in the dataset. Tab.7.6 presents the information from low resolution histograms together with the relative error on BV/TV on the resolved super resolution samples. As can be seen in the table, the super resolution images with highly underestimated BV/TV have one common character: the max values of the low resolution images are very huge compared with other samples, and their 'max/peak 2' are very high if they have one. Moreover, when the max values on the low resolution images are relatively low, the super resolution images potentially tends to have an overestimated BV/TV.

7.4 Conclusion and discussion

In this chapter, we proposed to use FSRCNN-Residual network to solve a joint super resolution/segmentation problem. The upsampling factors are 2 and 3.42. Most of the samples have been successfully recovered, while 4 samples lost a certain percents of bone mass. This may be related to the image statistics. To be precise, this may be related to the normalization of the data. The first results on the topological parameter Conn.D suggests that the preservation of Conn.D remains a challenge. In the following, we are going to discuss the solutions for these two drawbacks.

The way of normalization

For the normalization of the data, another choice is to normalize images before feeding them into the network: firstly, reduce the mean value of the image, then divided it by the standard deviation. This normalization way is expected to be better than simply normalizing the images via dividing by the max value, because it takes account of the entire distribution of image intensity.

How to preserve the connectivity?

Intuitively, there are two ideas to preserve the connectivity. The first one is to add a third term in the loss function which reflects neighbor information. The current loss function is based on the pixel-wise accuracy. The neighbor information increased the perception region, thus a larger local neighbor region can be considered.

The second idea is inspired from (Oktay et al., 2018). Both estimation and ground truth images can be compacted into a low dimension space reflecting their connectivity, then their compacted features are optimized to be similar so that the connectivity is preserved or recovered.

Chapter 8

Conclusions and Perspectives

Conclusions

In this thesis, we attempted to solve a joint super resolution/segmentation problem for experimental 3D HR-pQCT trabecular bone micro architecture images. Experimental HR-pQCT images of bone with a voxel size of $82\mu\text{m}$ and μCT images with a voxel size at $24\mu\text{m}$ were available. The μCT images were considered as the ground truth. Our goal was to improve the resolution of HR-pQCT images so that the bone parameters extracted from the segmented super resolution images are close to the ones obtained with μCT image measurements.

From chapter 3 to chapter 5, the upsampling factor of super resolution was 2. We registered the raw μCT images so that the voxel size turned to be $41\mu\text{m}$ instead of $24\mu\text{m}$. In chapter 7, we tried to solve a joint super resolution/segmentation problem with the upsampling factor 3.42, which means that the voxel size has changed from $82\mu\text{m}$ to $24\mu\text{m}$.

In the third chapter, we have used the notion of mutual information to determine the degradation kernel of HR-pQCT and compare it with the kernel determined with classic L_2 norm minimization. Precisely, we apply both kernels to solve the super resolution problem with TV regularization on experimental HR-pQCT images. Our numerical experiments show that the results obtained with both kernels are comparable. The resolved super resolution images present obvious artifacts and the contrast of the restored images is not as good as the one of the ground truth images.

In chapter 4, in order to recover the good contrast of the high resolution images, we have added a double-well nonconvex potential as regularization term. This is a nonconvex and nonsmooth optimization problem, and we have compared three schemes to minimize the regularization functional. The results obtained with three different approaches are similar. All of them have enhanced the contrast of the images, but there are many structural details that are lost.

In chapter 5, we have applied a 2.5D SCDL method to preserve the structural information. Specifically, two dictionaries are learned in the method: one for high resolution images, the other for low resolution images. Both dictionaries are associated with linear mapping functions. The resolved super resolution images are visually improved, but the connectivity density remains a challenge.

One difficulty in dictionary learning approach to solve super resolution problem is how to couple the two dictionaries properly. In the literature, different mapping functions are

discussed, and most of them use linear function to associates dictionaries of high and low resolution images. We believe deep learning based approaches are able to describe more complex mapping functions.

A review of the super resolution techniques in deep learning is given in chapter 6. The evolution of architectures in super resolution is presented. Meanwhile, the application of deep learning based super resolution in medical image processing is discussed.

In chapter 7, we have proposed a CNN based approach to solve the super resolution and segmentation problems at the same time. The network is inspired from two networks, one is FSRCNN that tackles super resolution problems, the other is the Residual network. The residual network has been proposed to solve super resolution problems, but in this work, we used it to refine the segmentation of images. The numerical experiments demonstrated that the images have been improved and the estimation of bone parameters becomes more accurate. In fact, the results obtained with the FSRCNN network and 2.5D SCDL are comparable. But FSRCNN-Residual network outperforms SCDL regarding time consumption. During the training process, the 2.5D SCDL updates dictionary atoms sequentially, which largely slows down the efficiency of method, while FSRCNN-Residual network is optimized with parallel calculation. Moreover, 2.5D SCDL for 3D image reconstruction requires to proceed super resolution method by slices in three directions, thus it is less efficient than FSRCNN-Residual network which outputs 3D super resolution volume directly.

In the original dataset, the voxel sizes of high resolution and low resolution images are $24\mu\text{m}$ and $82\mu\text{m}$. This is a super resolution problem with upsampling factor at 3.42. Modeling the problem as a super resolution with upsampling factor 2 is a compromise since super resolution with a factor 3.42 is not an easy problem to solve with variational methods. But we found that CNN made the super resolution problem of factor 3.42 possible to handle. Furthermore, our experimental results show that FSRCNN-Residual network is an optimal choice for the super resolution problem with upsampling factor 2. For these reasons, we have applied FSRCNN-Residual to deal with the super resolution problem with upsampling factor equal to 3.42. The parallel calculation with GPU permits to reconstruct super resolution images for entire bone samples.

However, one drawback of the current FSRCNN-Residual network is that when the distribution of image intensity is not uniform, the resolved images may loss many structural information. Moreover, how to preserve and recover the connectivity remains a challenge.

Perspectives

Improvement for the current network

The results presented in chapter 7 are preliminary work. The proposed network can be potentially improved by changing the way of normalization.

Moreover, as mentioned in chapter 6, the residual network may be replaced by a densely connected network. The densely connected network reduces the number of parameters in the network by sharing the features, and the performance is better compared to other architectures in the literature.

How to preserve the connectivity

One way that may preserve the connectivity is to introduce an additional loss that takes account of the local neighbor information. The current network has two losses, but both of them calculate the accuracy of the estimation element-wise. The constraints on neighbors similarity may boost the performance on connectivity.

Furthermore, inspired from the anatomically constrained neuron network, enhancing the similarity of the estimation and high resolution images in low dimension space may help to preserve the connectivity. The challenge is how to define the low dimension space which describes the topological structure.

List of publications

Journal

- [J1] Yufei LI, Bruno SIXOU, Françoise PEYRIN, "Nonconvex mixed TV/Cahn-Hilliard functional for super-resolution/segmentation of 3D Trabecular bone images", accepted in *Journal of Mathematical Imaging and Vision*, 10/2018.
- [J2] Yufei LI, Bruno SIXOU, Andrew BURGHARDT, Françoise PEYRIN, "Investigation of semi-coupled dictionary learning in 3D super resolution HR-pQCT imaging", submit in *IEEE Transactions on Radiation and Plasma Medical Sciences*, 11/2018.

Conference papers

- [C1] Bruno SIXOU, Yufei LI, Françoise PEYRIN, "Determination of blur kernel for HR-pQCT with bilevel optimization" NCMIP 2018
- [C2] Yufei LI, Bruno SIXOU, Françoise PEYRIN, "Estimation of the blurring kernel in experimental HR-pQCT images based on mutual information", EUSIPCO 2017.
- [C3] Yufei LI, Bruno SIXOU, Françoise PEYRIN, "Super resolution/segmentation of 3D trabecular bone images with total variation and nonconvex cahn-Hilliard functional", ISBI 2017.
- [C4] Yufei LI, Bruno SIXOU, Françoise PEYRIN, "Super-resolution/segmentation of 2D trabecular bone images by a Mumford-Shah approach and comparison to Total Variation", EUSIPCO 2016.

Articles under preparation

- [P1] Yufei LI, Bruno SIXOU, Matthieu MARTIN, Philippe DELACHARTRE, Andrew BURGHARDT, Françoise Peyrin, "solving a joint super resolution/segmentation problem for 3D HR-pQCT images via a convolutional neural network", *EUSIPCO 2019*.
- [P2] Yufei LI, Bruno Sixou, Françoise PEYRIN, "A deep learning method to assess human bone micro architecture from super resolved 3D HR-pQCT images", *JBMR*.
- [P3] Yufei LI, Bruno SIXOU, Françoise PEYRIN, "A review for the applications of deep learning in medical imaging super resolution" *IEEE Transaction Medical Imaging*.

Annexes

Optimization Preliminaries

Before the introduction of convex or nonconvex optimization, we first briefly review some preliminary notions (Artacho et al., 2014, Gasso et al., 2009). We consider a Banach space $\mathcal{X}$ and its dual $\mathcal{X}^*$.

Convex set. Let $\mathbb{C} \subseteq \mathcal{X}$, if $\forall \mathbf{f}, \mathbf{v} \in \mathbb{C}, \mu \in [0, 1]$

$$\mu \mathbf{f} + (1 - \mu) \mathbf{v} \in \mathbb{C} \quad (8.1)$$

then $\mathbb{C}$ is a convex set.

Indicator function.

$$\iota_{\mathbb{C}}(\mathbf{f}) = \begin{cases} 0, & \text{if } \mathbf{f} \in \mathbb{C}; \\ +\infty, & \text{otherwise.} \end{cases}$$

Convex function. Let $\mathcal{F} : \mathbb{C} \rightarrow]-\infty, +\infty]$ be a function, $\text{dom} \mathcal{F} := \mathcal{F}^{-1}(\mathbb{R})$ is the domain of $\mathcal{F}$. $\mathcal{F}$ is a convex function if $\forall \mathbf{f}, \mathbf{v} \in \text{dom} \mathcal{F}, \mu \in [0, 1]$, we have:

$$\mathcal{F}(\mu \mathbf{f} + (1 - \mu) \mathbf{v}) \leq \mu \mathcal{F}(\mathbf{f}) + (1 - \mu) \mathcal{F}(\mathbf{v}) \quad (8.2)$$

A convex function $\mathcal{F}$ is said to be ν -strongly convex if $\exists \nu > 0$, such that $\mathcal{F}(\mathbf{f}) - \frac{\nu}{2} \|\mathbf{f}\|^2$ is convex.

Subdifferential. If $\text{dom} \mathcal{F} \neq \emptyset$, $\mathcal{F}$ is a proper function. For proper function $\mathcal{F}$, $\mathbf{f}^* \in \mathcal{X}^*$, its subdifferential is defined by

$$\partial \mathcal{F} : \mathcal{X} \rightarrow \mathcal{X}^* : \mathbf{f} \mapsto \{\mathbf{f}^* \in \mathcal{X}^* : \mathcal{F}(\mathbf{v}) \geq \mathcal{F}(\mathbf{f}) + \langle \mathbf{v} - \mathbf{f}, \mathbf{f}^* \rangle, \quad \forall \mathbf{v} \in \text{dom} \mathcal{F}\} \quad (8.3)$$

One important property of convex function is that its subdifferential is a monotone operator.

Lower-semicontinuity. $\mathcal{F} : \mathcal{X} \rightarrow]-\infty, +\infty]$, $\mathcal{F}$ is lower-semicontinuous if

$$\liminf_{\mathbf{f} \rightarrow \bar{\mathbf{f}}} \mathcal{F}(\mathbf{f}) \geq \mathcal{F}(\bar{\mathbf{f}}), \quad \forall \bar{\mathbf{f}} \in \mathcal{X}. \quad (8.4)$$

Fenchel conjugate. Let $\mathcal{F} : \mathcal{X} \rightarrow]-\infty, +\infty]$, for $\mathbf{f}^* \in \mathcal{X}^*$ and $\mathbf{f} \in \mathcal{X}$, we denote $\mathcal{F}^*$ the Fenchel conjugate function of $\mathcal{F}$,

$$\mathcal{F}^*(\mathbf{f}^*) = \sup_{\mathbf{f} \in \mathcal{X}} \{\langle \mathbf{f}, \mathbf{f}^* \rangle - \mathcal{F}(\mathbf{f})\} \quad (8.5)$$

Fenchel-Young inequality. Let $\mathcal{F} : \mathcal{X} \rightarrow]-\infty, +\infty]$, for $\mathbf{f}^* \in \mathcal{X}^*$ and $\mathbf{f} \in \mathcal{X}$,

$$\mathcal{F}^*(\mathbf{f}^*) + \mathcal{F}(\mathbf{f}) \geq \langle \mathbf{f}^*, \mathbf{f} \rangle. \quad (8.6)$$

Equality holds if and only if $\mathbf{f}^* \in \partial \mathcal{F}(\mathbf{f})$

Fenchel-Young inequality is very useful for convex and nonconvex optimization problems. When a proper low-semicontinuous function $\mathcal{F}$ is convex, its biconjugate function $\mathcal{F}^{**}$ is itself $\mathcal{F}$ (Fenchel–Moreau theorem). When $\mathcal{F}$ is not convex, its biconjugate function is the convex hull of $\mathcal{F}$.

Registration

One note for the image order: since each .ISO file of HR-pQCT contains 5 bone samples, the first column the table below presents sources of tomography images.

HR-pQCT	C0008233.ISO	C0008232.ISO	C0008231.ISO	C0008230.ISO
μ CT	D0000812.ISO	D0000813.ISO	D0000814.ISO	D0000818.ISO
μ CT	E0000812.ISO	E0000813.ISO	E0000814.ISO	E0000818.ISO
μ CT	F0000812.ISO	F0000813.ISO	F0000814.ISO	F0000818.ISO
μ CT	G0000812.ISO	G0000813.ISO	G0000814.ISO	G0000818.ISO
μ CT	H0000812.ISO	H0000813.ISO	H0000814.ISO	H0000818.ISO

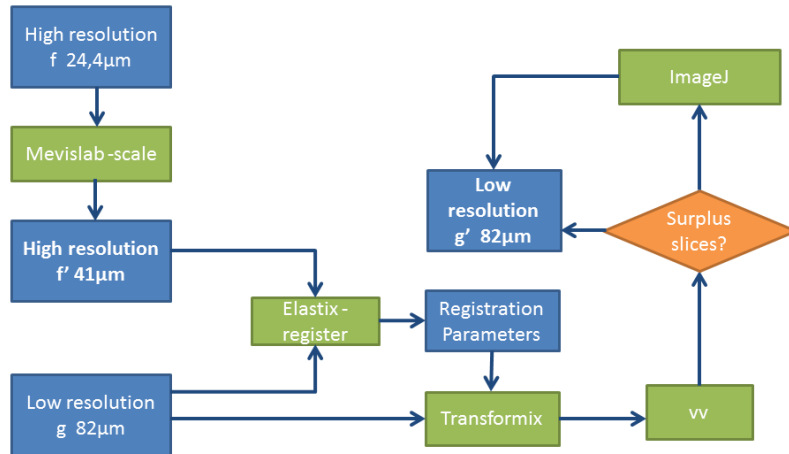


Figure 8.1: Diagram of registration

There are two resolutions available, one is $24.4\mu m$, denoted by $\mathbf{f}$, the other is $82\mu m$, denoted by $\mathbf{g}$. How to do the register when the resolution of one volume is not the integer multiples of the other one? I followed the work of Alina TOMA. For the preparing

work, we need to use 'ImageJ' to crop the image into a proper size, eliminate unnecessary black background to reduce the complexity of computation, install Mevislab, vv as well as elastix.. Firstly, $\mathbf{f}$ is downsampled to $41\mu m$ by Mevislab with Resampling model, the output is noted by $\mathbf{f}'$. Then we regard $\mathbf{f}'$ to be the ground truth of the $\mathbf{g}$, and start to do registration. Basic idea is to register $\mathbf{f}'$ with $\mathbf{g}$. By using the following command:

```
elastix -f fixedImage.ext -m movingImage.ext
-out outputDirectory -p parameterFile.txt
```

Our .ext files are .mhd images. The fixed image mentioned in the command is $\mathbf{f}'$; the moving image is $\mathbf{g}$. If we note $\mathbf{g}_r$ to be the output of this command, then the resolution of $\mathbf{g}_r$ is $42\mu m$. Our work aims to improve the CT images from HR-pQCT, which is regarded as the low resolution image, then we need to find a $\mathbf{g}'$ with the same resolution as $\mathbf{g}$, being in the same rotation and offsets as $\mathbf{g}_r$ without scaling. There is an output file named "TransformParameters.0.txt", which offers registration information like rotation and stretching parameters. Then we could use transformix to solve this problem. Firstly, copy "TransformParameters.0.txt" and renamed as 'TransformParameters.txt'. Next change the following parameters in the "TransformParameters.txt" to be:

```
(Size 576 447 144)
(Spacing 0.082 0.082 0.082)
(Origin 0 0 0)
```

These values are supposed to be consistent with the low resolution image $\mathbf{g}$. Then call transformix program in console with the following command:

```
transformix -in inputImage.ext -out outputDirectory
-tp TransformParameters.txt
```

Open the output .mhd image and $\mathbf{f}'$ to see if there are some extra slices. Eliminate surplus slices if there are some. Finally the output will be a good $\mathbf{g}'$, with $82\mu m$ resolution, well registered with respect to $\mathbf{f}'$.

New Registration (only move one sample)

In the former registration work, both high and low resolution images have been changed, thus the structure information is modified as well. Take sample pairs C0008230-D0000818 as an example, their structural parameters are listed in the table 8.1

Even though such registration procedure changes structure information, it is still a good solution when the upsampling factor of images is not integer. Thanks to deep learning frame work, we can get rid of such constraint: even if the upsampling factor is not integer. So in this section, we describe how to perform registration between sample at voxel size $82\mu m$ and $24\mu m$. Let sample at voxel size $82\mu m$ and $24\mu m$ be $\mathbf{g}$ and $\mathbf{f}_{24}$. In principle, the idea is to down sample the high resolution image to $82\mu m$ and proceed registration to the low resolution image. The obtained transformation parameters are then applied to the original high resolution image. The details information are given in the following.

Table 8.1: Different resolution images are involved in a sample registration process. Here we analyse sample 'C0008230-D0000818' and the table summarized topological criteria of different resolution images. This table only considered Bone volume and Connectivity, but not BV/TV or Conn.D, since BV/TV or Conn.D are computed based on total volume of the sample. For a fixed sample at different resolution, firstly, they share the sample total volume, and secondly, if we generate mask for different resolution images to obtain total volume corresponding to different resolution, errors may be introduced.

parameter\samples	D0000818 (orginal)	D0000818-41 (sampled)	C0008230 (orginal)	C0008230 (registered)
voxel size	24	41	82	82
Component Nbr	101	27	138	219
Cavities Nbr	1485	166	44	26
BV	835.73	907.26	1519.13	1105.92
Connectivity	7542.00	7185.00	9894.00	6035.00
Bone surface	6979.48	6614.29	8886.67	6077.20
SMI	1.03	1.56	2.69	3.87

downsampling

Start from an original high resolution image (f_{24}), with voxel size equal to $24\mu\text{m}$. We use MeVisLab to down sample f_{24} to voxel size $81\mu\text{m}$. The down sampling operator is cubic B-spline. Here we denote the output as f_{82} .

registration

The tool for registration is the same as the one in the previous work: elastix. This is an open source software for registration, it can be integrated with Matlab or ImageJ. The command for registration is

```
elastix.exe -f g.mhd -m f_82.mhd -out outputPath
-p newElastix.txt -fMask g-mask.mhd
```

This command means that in this registration, the fixed volume is 'g.mhd' (low resolution image), the moved one is f_{82} , with registration parameters defined in 'newElastix.txt' and a mask of 'g.mhd'. Even though conventional registration keeps high resolution image fixed while moving low resolution image, such operation may results in low resolution structural information changes. If we use a modified low resolution image to solve super resolution problem, even though finally we have a good results, it may not be fit practical problems. After all, we can't know how to transform the original data to the desired inputs without reference images.

In the registration process, the mask here not only saves computation cost, but also removes unnecessary or even negative contribution of image background. In elastix, both fixed and moved image masks can be provided, but it is suggested that only the fixed image mask may be sufficient. In this work, the mask of the fixed volume is generated by 4 times dilation with sphere of radius equal to 5. Then fill in the holes within the mask. Not necessary to attribute labels to image objects. For instance, in 1D, a line starts from

A, ends at B, if this line is broken, then broken ends named as C and D. If CD is missing, just connecting A and B can fill in such 'hole', the same principle in 3D.

As for registration parameters, here is the code:

```
//Image types
(FixedInternalImagePixelType "short")
(MovingInternalImagePixelType "short")

//Image dimensions
(FixedImageDimension 3)
(MovingImageDimension 3)
(UseDirectionCosines "true")

// ***** Main Components *****
(Registration "MultiResolutionRegistration")
(NumberOfResolutions 1)
(FixedImagePyramid "FixedSmoothingImagePyramid")
(MovingImagePyramid "MovingSmoothingImagePyramid")
//(Interpolator "Interpolator")
(Optimizer "AdaptiveStochasticGradientDescent")
(Transform "EulerTransform")
(Metric "AdvancedMattesMutualInformation")
//(Metric "AdvancedMeanSquares")

// ***** Transformation *****
//Compose the transformations
//(HowToCombineTransforms "Compose")
(AutomaticTransformInitialization "true")
(AutomaticScalesEstimation "true")
(AutomaticTransformInitializationMethod "CenterOfGravity")

// ***** Optimizer *****
(MaximumNumberOfIterations 300)

// ***** Image sampling *****
//Random sampler over the voxels of the image
(ImageSampler "RandomCoordinate")
(NumberOfSpatialSamples 5000)
(NewSamplesEveryIteration "true")

// ***** Metric *****
//(NumberOfHistogramBins 32)

//*****Several*****
(ErodeMask "true")

// ***** Interpolation and Resampling *****
```

```
//Default pixel value for pixels that come from
//outside the picture:
(DefaultPixelValue 0)
//Write or no the images
(WriteResultImage "true")
// The pixel type and format of the resulting
//deformed moving image
(ResultImagePixelFormat "short")
(ResultImageFormat "mhd")
```

The output of the above command not only includes registered volumes, but also a file named as "TransformParameters.0.txt", including relative transform parameters, such as rotation angles, translation parameters and so on. This file will be used in the follow-up transformation process.

Before moving on to next step, it would be better to visually verify if images are well registered. We use vv.exe to observe images.

Transformation

Elastix allows to transform images with a set of given parameters, thereafter we can use "TransformParameters.0.txt" to transform the original high resolution image. Still some parameters to modify, such as element space and output size. The output size can be initialized as the size of original high resolution image, but it is completely possible that, in the output file, only part of sample appears, and other slices may get out of the range of the preset output size after rotation and translation. So it would be better to set the output size be 3 or 4 times of low resolution images. But visualization is always necessary, even though such configuration usually defines an output region enough large, but in some special cases, it could be still insufficient.

In this step, "overlay mode" in vv.exe allows to visualize the effects of registration between the new output image (voxel size $24\mu\text{m}$) and the low resolution image(voxel size $82\mu\text{m}$). If all slices are involved in the registration, you will observe that last slices in the stack are zeros or no borders are cut abruptly,

In addition, it will be even better to compare topological parameters of original and transformed high resolution images (voxel size $24\mu\text{m}$) to see if the transformed image is not too far away from the original one. Please note that if the size of volumes are different, even though both images describe the same bone structure, additional zeros influences segmentation results, thus their parameters are not similar.

In order to solve this problem, we consider to use one prior which is in the original high resolution CT images, the intensity of voxels belonging to background are non-zeros, in another word, zeros appear rarely, thereafter the negative effects of zero padding can be removed by the Otsu segmentation with Matlab command

```
seg=zeros(size(img));
% img~=0 can be regarded as index of non zero voxels.
temp=img(img~=0);
seg(img~=0)=mat2gray(temp)>graythresh(mat2gray(temp))
```

parameter\samples	D0000818 (orginal)	D0000818-41 (registered)	C0008230 (orginal)
voxel size	24	24	82
Component Nbr	101	104	138
Cavities Nbr	1485	1385	44
BV	835.73	888.86	1519.13
Connectivity	7542.00	7307.00	9894.00
Bone surface	6979.48	7508.73	8886.67
SMI	1.03	0.95	2.69

Table 8.2: Parameter synthesis after new registration

Table ?? presents topological parameter comparison in new registration procedure. We could see that even though the registered sample (at third column) is not exactly the same as, but quite approach to the original high resolution image. If the new transformed image not only has a similar parameter compared to the original high resolution image (voxel size $24\mu\text{m}$), but also well overlaps with low resolution image (voxel size $82\mu\text{m}$), then the new registration is proved to be reliable. In latter in our super resolution work, it can be regarded as the reference image. However, we still have another problem to handle before moving on to solve super resolution problem: how to trim the reference and low resolution images so that their size are properly fit, especially when the upsampling factor is not integer.

Volume trimming

41 voxels with voxel space equal to $24\mu\text{m}$ has the same length as 12 voxels with voxel size equal to $41\mu\text{m}$. With this observation, we grid high and low segmented resolution image with interval 41 and 12 voxels respectively. Since the output size of transform is not matched with low resolution image, and usually we want to ensure all slices are involved in registration, the transformed output (high resolution image) usually have more grids than low resolution images'. Under the assumption that both images starts from the same origin, in order to trim both images properly, here we propose to preserve the grids sharing the same index.

Bibliography

- Abadi, M., Barham, P., Chen, J., Chen, Z., Davis, A., Dean, J., Devin, M., Ghemawat, S., Irving, G., Isard, M., et al. (2016). Tensorflow: A system for large-scale machine learning. In *OSDI*, volume 16, pages 265–283.
- Aharon, M., Elad, M., and Bruckstein, A. (2006). *rmk-svd*: An algorithm for designing over-complete dictionaries for sparse representation. *IEEE Transactions on signal processing*, 54(11):4311–4322.
- Anbarjafari, G. and Demirel, H. (2010). Image super resolution based on interpolation of wavelet domain high frequency subbands and the spatial domain input image. *ETRI journal*, 32(3):390–394.
- Artacho, F. J. A., Borwein, J. M., Martín-Márquez, V., and Yao, L. (2014). Applications of convex analysis within mathematics. *Mathematical Programming*, 148(1-2):49–88.
- Artina, M., Fornasier, M., and Solombrino, F. (2013). Linearly constrained nonsmooth and nonconvex minimization. *SIAM Journal on Optimization*, 23(3):1904–1937.
- Asif, M., Akram, M. U., Hassan, T., Shaukat, A., and Waqar, R. (2017). High resolution OCT image generation using super resolution via sparse representation. In *Eighth International Conference on Graphic and Image Processing (ICGIP 2016)*, volume 10225, page 1022512. International Society for Optics and Photonics.
- Asif, M., Khan, S. A., Hassan, T., Akram, M. U., and Shaukat, A. (2016). Generation of High Resolution Medical Images Using Super Resolution via Sparse Representation. In *International Afro-European Conference for Industrial Advancement*, pages 288–298. Springer.
- Bertozzi, A., Esedoğlu, S., and Gillette, A. (2007a). Analysis of a two-scale Cahn-Hilliard model for binary image inpainting. *Multiscale Modeling & Simulation*, 6(3):913–936.
- Bertozzi, A. L., Esedoğlu, S., and Gillette, A. (2007b). Inpainting of binary images using the Cahn-Hilliard equation. *IEEE Transactions on image processing*, 16(1):285–291.
- Bevilacqua, M., Roumy, A., Guillemot, C., and Alberi-Morel, M. L. (2012). Low-complexity single-image super-resolution based on nonnegative neighbor embedding.
- Bolotin, H. (2007). DXA in vivo BMD methodology: an erroneous and misleading research and clinical gauge of bone mineral status, bone fragility, and bone remodelling. *Bone*, 41(1):138–154.
- Bolte, J., Sabach, S., and Teboulle, M. (2014). Proximal alternating linearized minimization or nonconvex and nonsmooth problems. *Mathematical Programming*, 146(1-2):459–494.

- Bonnassie, A., Peyrin, F., and Attali, D. (2003). A new method for analyzing local shape in three-dimensional images based on medial axis transformation. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 33(4):700–705.
- Borman, S. and Stevenson, R. (1998). Spatial resolution enhancement of low-resolution image sequences - a comprehensive review with directions for future research. *Lab. Image and Signal Analysis, University of Notre Dame, Tech. Rep.*
- Bosq, D. (2012). *Nonparametric statistics for stochastic processes: estimation and prediction*, volume 110. Springer Science & Business Media.
- Boutroy, S., Bouxsein, M. L., Munoz, F., and Delmas, P. D. (2005). In vivo assessment of trabecular bone microarchitecture by high-resolution peripheral quantitative computed tomography. *The Journal of Clinical Endocrinology & Metabolism*, 90(12):6508–6515.
- Boyd, S. (2011). Alternating direction method of multipliers. In *Talk at NIPS Workshop on Optimization and Machine Learning*.
- Boyd, S. and Vandenberghe, L. (2004). *Convex optimization*. Cambridge university press.
- Brandi, M. L. (2009). Microarchitecture, the key to bone quality. *Rheumatology*, 48(suppl_4):iv3–iv8.
- Burghardt, A. J., Buie, H. R., Laib, A., Majumdar, S., and Boyd, S. K. (2010). Reproducibility of direct quantitative measures of cortical bone microarchitecture of the distal radius and tibia by HR-pQCT. *Bone*, 47(3):519–528.
- Burghardt, A. J., Link, T. M., and Majumdar, S. (2011). High-resolution computed tomography for clinical imaging of bone microarchitecture. *Clinical Orthopaedics and Related Research*, 469(8):2179–2193.
- Burghardt, A. J., Pialat, J.-B., Kazakia, G. J., Boutroy, S., Engelke, K., Patsch, J. M., Valentinitsch, A., Liu, D., Szabo, E., Bogado, C. E., et al. (2013). Multicenter precision of cortical and trabecular bone quality measures assessed by high-resolution peripheral quantitative computed tomography. *Journal of Bone and Mineral Research*, 28(3):524–536.
- Campion, J. M. and Maricic, M. J. (2003). Osteoporosis in men. *American family physician*, 67(7):1521–1526.
- Candes, E. J., Wakin, M. B., and Boyd, S. P. (2008). Enhancing sparsity by reweighted l_1 minimization. *Journal of Fourier analysis and applications*, 14(5-6):877–905.
- Cannon, M. (1976). Blind deconvolution of spatially invariant image blurs with phase. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 24(1):58–63.
- Cao, T., Singh, N., Jovic, V., and Niethammer, M. (2015). Semi-coupled dictionary learning for deformation prediction. In *Biomedical Imaging (ISBI), 2015 IEEE 12th International Symposium on*, pages 691–694. IEEE.
- Cao, T., Zach, C., Modla, S., Powell, D., Czymmek, K., and Niethammer, M. (2014). Multi-modal registration for correlative microscopy using image analogies. *Medical image analysis*, 18(6):914–926.
- Carey, W. K., Chuang, D. B., and Hemami, S. S. (1999). Regularity-preserving image interpolation. *IEEE transactions on image processing*, 8(9):1293–1297.

- Carpintero, P., Caeiro, J. R., Carpintero, R., Morales, A., Silva, S., and Mesa, M. (2014). Complications of hip fractures: A review. *World journal of orthopedics*, 5(4):402.
- Chambolle, A., Caselles, V., Cremers, D., Novaga, M., and Pock, T. (2010). An introduction to total variation for image analysis. *Theoretical foundations and numerical methods for sparse recovery*, 9(263-340):227.
- Chang, M. M., Tekalp, A. M., and Erdem, A. T. (1991). Blur identification using the bispectrum. *IEEE transactions on signal processing*, 39(10):2323–2325.
- Chen, T., Li, M., Li, Y., Lin, M., Wang, N., Wang, M., Xiao, T., Xu, B., Zhang, C., and Zhang, Z. (2015). Mxnet: A flexible and efficient machine learning library for heterogeneous distributed systems. *arXiv preprint arXiv:1512.01274*.
- Chen, Y., Ranftl, R., and Pock, T. (2014). Insights into analysis operator learning: From patch-based sparse models to higher order mrfs. *IEEE Transactions on Image Processing*, 23(3):1060–1072.
- Chen, Y., Xie, Y., Zhou, Z., Shi, F., Christodoulou, A. G., and Li, D. (2018). Brain mri super resolution using 3D deep densely connected neural networks. In *Biomedical Imaging (ISBI 2018), 2018 IEEE 15th International Symposium on*, pages 739–742. IEEE.
- Chen, Y., Yin, X., Shi, L., Shu, H., Luo, L., Coatrieux, J.-L., and Toumoulin, C. (2013). Improving abdomen tumor low-dose CT images using a fast dictionary learning based processing. *Physics in Medicine & Biology*, 58(16):5803.
- Cheung, A. M., Adachi, J. D., Hanley, D. A., Kendler, D. L., Davison, K. S., Josse, R., Brown, J. P., Ste-Marie, L.-G., Kremer, R., Erlandson, M. C., et al. (2013). High-resolution peripheral quantitative computed tomography for the assessment of bone strength and structure: a review by the canadian bone strength working group. *Current osteoporosis reports*, 11(2):136–146.
- Combettes, P. L. and Pesquet, J.-C. (2011). Proximal splitting methods in signal processing. In *Fixed-point algorithms for inverse problems in science and engineering*, pages 185–212. Springer.
- Crouse, M. S., Nowak, R. D., and Baraniuk, R. G. (1998). Wavelet-based statistical signal processing using hidden markov models. *IEEE Transactions on signal processing*, 46(4):886–902.
- Denis, L., Thiébaud, E., Soulez, F., Becker, J.-M., and Mourya, R. (2015). Fast approximations of shift-variant blur. *International Journal of Computer Vision*, 115(3):253–278.
- Dice, L. R. (1945). Measures of the amount of ecologic association between species. *Ecology*, 26(3):297–302.
- Dong, C., Loy, C. C., He, K., and Tang, X. (2014). Learning a deep convolutional network for image super-resolution. In *European Conference on Computer Vision*, pages 184–199. Springer.
- Dong, C., Loy, C. C., and Tang, X. (2016). Accelerating the super-resolution convolutional neural network. In *European Conference on Computer Vision*, pages 391–407. Springer.

- Elleuch, I., Abdelkefi, F., Siala, M., Hamila, R., and Al-Dhahir, N. (2016). Quasi-sparsest solutions for quantized compressed sensing by graduated-non-convexity based reweighted l_1 minimization. In *Signal Processing Conference (EUSIPCO), 2016 24th European*, pages 473–477. IEEE.
- Elliott, C. M. (1989). The Cahn-Hilliard model for the kinetics of phase separation. In *Mathematical models for phase change problems*, pages 35–73. Springer.
- Engl, H. W., Hanke, M., and Neubauer, A. (1996). *Regularization of inverse problems*, volume 375. Springer Science & Business Media.
- Eyre, D. J. (1998). An unconditionally stable one-step scheme for gradient systems. *Unpublished article*.
- Favati, P., Lotti, G., Menchi, O., and Romani, F. (2010). Performance analysis of maximum likelihood methods for regularization problems with nonnegativity constraints. *Inverse Problems*, 26(8):085013.
- Feldkamp, L., Davis, L., and Kress, J. (1984). Practical cone-beam algorithm. *JOSA A*, 1(6):612–619.
- Fergus, R., Singh, B., Hertzmann, A., Roweis, S. T., and Freeman, W. T. (2006). Removing camera shake from a single photograph. In *ACM Transactions on Graphics (TOG)*, volume 25, pages 787–794. ACM.
- Figueiredo, C. P., Kleyer, A., Simon, D., Stemmler, F., d’Oliveira, I., Weissenfels, A., Museyko, O., Friedberger, A., Hueber, A. J., Haschka, J., et al. (2017). Methods for segmentation of rheumatoid arthritis bone erosions in high-resolution peripheral quantitative computed tomography (HR-pQCT). In *Seminars in arthritis and rheumatism*. Elsevier.
- Foucart, S. and Lai, M.-J. (2009). Sparsest solutions of underdetermined linear systems via l_q -minimization for $0 < q \leq 1$. *Applied and Computational Harmonic Analysis*, 26(3):395–407.
- Gabay, D. and Mercier, B. (1976). A dual algorithm for the solution of nonlinear variational problems via finite element approximation. *Computers & Mathematics with Applications*, 2(1):17–40.
- Gashti, M. P., Alimohammadi, F., Hulliger, J., Burgener, M., Oulevey-Aboulfadl, H., and Bowlin, G. L. (2012). Microscopic methods to study the structure of scaffolds in bone tissue engineering: a brief review. *Editor: A. Méndez-Vilas, Current Microscopy Contributions to Advances in Science and Technology*, 1:625–638.
- Gasso, G., Rakotomamonjy, A., and Canu, S. (2009). Recovering sparse signals with a certain family of nonconvex penalties and DC programming. *IEEE Transactions on Signal Processing*, 57(12):4686–4698.
- Gensheng, Z. (2010). *Medical image reconstruction, Chinese version*. Higher Education Press.
- Geusens, P., Chapurlat, R., Schett, G., Ghasem-Zadeh, A., Seeman, E., De Jong, J., and Van Den Bergh, J. (2014). High-resolution in vivo imaging of bone and joints: a window to microarchitecture. *Nature Reviews Rheumatology*, 10(5):304–313.
- Gluckman, J. (2003). Gradient field distributions for the registration of images. In *Image Processing, 2003. ICIP 2003. Proceedings. 2003 International Conference on*, volume 2, pages II–691. IEEE.

- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., and Bengio, Y. (2014). Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680.
- Gu, S., Zuo, W., Xie, Q., Meng, D., Feng, X., and Zhang, L. (2015). Convolutional sparse coding for image super-resolution. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 1823–1831.
- Hanocka, R. and Kiryati, N. (2015). Progressive blind deconvolution. In *International Conference on Computer Analysis of Images and Patterns*, pages 313–325. Springer.
- Hansen, P. C. (1999). The L-curve and its use in the numerical treatment of inverse problems.
- Harvey, N., Dennison, E., and Cooper, C. (2010). Osteoporosis: impact on health and economics. *Nature Reviews Rheumatology*, 6(2):99–105.
- He, L., Qi, H., and Zaretzki, R. (2013). Beta process joint dictionary learning for coupled feature spaces with application to single image super-resolution. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 345–352.
- Hermosillo, G., Ched'Hotel, C., and Faugeras, O. (2002). Variational methods for multimodal image matching. *International Journal of Computer Vision*, 50(3):329–343.
- Hildebrand, T. and RÜEGSEGG, P. (1997). Quantification of bone microarchitecture with the structure model index. *Computer Methods in Biomechanics and Bio Medical Engineering*, 1(1):15–23.
- Hollenbach, K. A., Barrett-Connor, E., Edelstein, S. L., and Holbrook, T. (1993). Cigarette smoking and bone mineral density in older men and women. *American journal of public health*, 83(9):1265–1270.
- Huang, G., Liu, Z., Van Der Maaten, L., and Weinberger, K. Q. (2017). Densely connected convolutional networks. In *CVPR*, volume 1, page 3.
- Huang, Y., Paisley, J., Lin, Q., Ding, X., Fu, X., and Zhang, X.-P. (2014). Bayesian non-parametric dictionary learning for compressed sensing MRI. *IEEE Transactions on Image Processing*, 23(12):5007–5019.
- Ioffe, S. and Szegedy, C. (2015). Batch normalization: Accelerating deep network training by reducing internal covariate shift. *arXiv preprint arXiv:1502.03167*.
- Ji, H. and Fermüller, C. (2006). Wavelet-based super-resolution reconstruction: theory and algorithm. In *European Conference on Computer Vision*, pages 295–307. Springer.
- Jia, Y., Shelhamer, E., Donahue, J., Karayev, S., Long, J., Girshick, R., Guadarrama, S., and Darrell, T. (2014). Caffe: Convolutional architecture for fast feature embedding. In *Proceedings of the 22nd ACM international conference on Multimedia*, pages 675–678. ACM.
- Jing, X.-Y., Zhu, X., Wu, F., You, X., Liu, Q., Yue, D., Hu, R., and Xu, B. (2015). Super-resolution person re-identification with semi-coupled low-rank discriminant dictionary learning. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 695–704.
- Johnell, O. and Kanis, J. (2006). An estimate of the worldwide prevalence and disability associated with osteoporotic fractures. *Osteoporosis international*, 17(12):1726–1733.

- Joshi, N., Szeliski, R., and Kriegman, D. J. (2008). PSF estimation using sharp edge prediction. In *Computer Vision and Pattern Recognition, 2008. CVPR 2008. IEEE Conference on*, pages 1–8. IEEE.
- Kabanikhin, S. I. (2008). Definitions and examples of inverse and ill-posed problems. *Journal of Inverse and Ill-Posed Problems*, 16(4):317–357.
- Kamel, H. K., Hussain, M. S., Tariq, S., Perry III, H. M., and Morley, J. E. (2000). Failure to diagnose and treat osteoporosis in elderly patients hospitalized with hip fracture. *The American journal of medicine*, 109(4):326–328.
- Karimi, D. and Ward, R. (2016). Reducing streak artifacts in computed tomography via sparse representation in coupled dictionaries. *Medical physics*, 43(3):1473–1486.
- Kaufman, J.-M., Reginster, J.-Y., Boonen, S., Brandi, M., Cooper, C., Dere, W., Devogelaer, J.-P., Diez-Perez, A., Kanis, J. A., McCloskey, E., et al. (2013). Treatment of osteoporosis in men. *Bone*, 53(1):134–144.
- Khosla, S., Amin, S., and Orwoll, E. (2008). Osteoporosis in men. *Endocrine reviews*, 29(4):441–464.
- Khosla, S., Riggs, B. L., Atkinson, E. J., Oberg, A. L., McDaniel, L. J., Holets, M., Peterson, J. M., and Melton, L. J. (2006). Effects of sex and age on bone microstructure at the ultradistal radius: a population-based noninvasive in vivo assessment. *Journal of Bone and Mineral Research*, 21(1):124–131.
- Kim, J., Kwon Lee, J., and Mu Lee, K. (2016a). Accurate image super-resolution using very deep convolutional networks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1646–1654.
- Kim, J., Kwon Lee, J., and Mu Lee, K. (2016b). Deeply-recursive convolutional network for image super-resolution. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1637–1645.
- Kim, S. and Su, W.-Y. (1993). Subpixel accuracy image registration by spectrum cancellation. In *Acoustics, Speech, and Signal Processing, 1993. ICASSP-93., 1993 IEEE International Conference on*, volume 5, pages 153–156. IEEE.
- Kingma, D. P. and Ba, J. (2014). Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*.
- KK, N., ML, B., et al. (2017). Effects of Denosumab and Teriparatide Transitions on Bone Microarchitecture and Estimated Strength: the DATA-Switch HR-pQCT study. *Journal of Bone and Mineral Research*.
- Klein, S., Staring, M., Murphy, K., Viergever, M. A., and Pluim, J. P. (2010). Elastix: a tool-box for intensity-based medical image registration. *IEEE transactions on medical imaging*, 29(1):196–205.
- Kotera, J., Šroubek, F., and Milanfar, P. (2013). Blind deconvolution using alternating maximum a posteriori estimation with heavy-tailed priors. In *International Conference on Computer Analysis of Images and Patterns*, pages 59–66. Springer.

- Krizhevsky, A., Sutskever, I., and Hinton, G. E. (2012). Imagenet classification with deep convolutional neural networks. In *Advances in neural information processing systems*, pages 1097–1105.
- Lai, W.-S., Huang, J.-B., Ahuja, N., and Yang, M.-H. (2017). Deep laplacian pyramid networks for fast and accurate super resolution. In *IEEE Conference on Computer Vision and Pattern Recognition*, volume 2, page 5.
- LeCun, Y., Bengio, Y., and Hinton, G. (2015). Deep learning. *nature*, 521(7553):436.
- Ledig, C., Theis, L., Huszár, F., Caballero, J., Cunningham, A., Acosta, A., Aitken, A., Tejani, A., Totz, J., Wang, Z., et al. (2016). Photo-realistic single image super-resolution using a generative adversarial network. *arXiv preprint*.
- Li, H. and Lin, Z. (2015). Accelerated proximal gradient methods for nonconvex programming. In *Advances in Neural Information Processing Systems*, pages 379–387.
- Li, S., Yin, H., and Fang, L. (2012). Group-sparse representation with dictionary learning for medical image denoising and fusion. *IEEE Transactions on biomedical engineering*, 59(12):3450–3459.
- Li, Y., Sixou, B., Burghardt, A., and Peyrin, F. (2017a). Super-resolution/segmentation of 3D trabecular bone images with total variation and nonconvex Cahn-Hilliard functional. In *Biomedical Imaging (ISBI 2017), 2017 IEEE 14th International Symposium on*, pages 1193–1196. IEEE.
- Li, Y., Sixou, B., and Peyrin, F. (2017b). Estimation of the blurring kernel in experimental HR-pQCT images based on mutual information. In *Signal Processing Conference (EUSIPCO), 2017 25th European*, pages 2086–2090. IEEE.
- Lim, B., Son, S., Kim, H., Nah, S., and Lee, K. M. (2017). Enhanced deep residual networks for single image super-resolution. In *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, volume 1, page 3.
- Litjens, G., Kooi, T., Bejnordi, B. E., Setio, A. A. A., Ciompi, F., Ghafoorian, M., van der Laak, J. A., Van Ginneken, B., and Sánchez, C. I. (2017). A survey on deep learning in medical image analysis. *Medical image analysis*, 42:60–88.
- Liu, X. S., Zhang, X. H., Sekhon, K. K., Adams, M. F., McMahon, D. J., Bilezikian, J. P., Shane, E., and Guo, X. E. (2010). High-resolution peripheral quantitative computed tomography can assess microstructural and mechanical properties of human distal tibial bone. *Journal of Bone and Mineral Research*, 25(4):746–756.
- Lucchese, L. and Cortelazzo, G. M. (2000). A noise-robust frequency domain technique for estimating planar roto-translations. *IEEE Transactions on Signal Processing*, 48(6):1769–1786.
- Luo, J., Eri, H., Can, A., Ramani, S., Fu, L., and De Man, B. (2016). 2.5D dictionary learning based computed tomography reconstruction. In *SPIE Defense+ Security*, pages 98470L–98470L. International Society for Optics and Photonics.
- MacNeil, J. A. and Boyd, S. K. (2007). Accuracy of high-resolution peripheral quantitative computed tomography for measurement of bone quality. *Medical engineering & physics*, 29(10):1096–1105.

- Mairal, J., Bach, F., Ponce, J., and Sapiro, G. (2009). Online dictionary learning for sparse coding. In *Proceedings of the 26th annual international conference on machine learning*, pages 689–696. ACM.
- Mairal, J., Bach, F., Ponce, J., and Sapiro, G. (2010). Online learning for matrix factorization and sparse coding. *Journal of Machine Learning Research*, 11(Jan):19–60.
- Manske, S. L., Zhu, Y., Sandino, C., and Boyd, S. K. (2015). Human trabecular bone microarchitecture can be assessed independently of density with second generation HR-pQCT. *Bone*, 79:213–221.
- Mansoor, A., Vongkovit, T., and Linguraru, M. G. (2018). Adversarial approach to diagnostic quality volumetric image enhancement. In *Biomedical Imaging (ISBI 2018), 2018 IEEE 15th International Symposium on*, pages 353–356. IEEE.
- Marcel, B., Briot, M., and Murrieta, R. (1997). Calcul de translation et rotation par la transformation de Fourier. *TS. Traitement du signal*, 14(2):135–149.
- Martín-Badosa, E., Amblard, D., Nuzzo, S., Elmoutaouakkil, A., Vico, L., and Peyrin, F. (2003). Excised bone structures in mice: imaging at three-dimensional synchrotron radiation micro CT. *Radiology*, 229(3):921–928.
- McCann, M. T., Jin, K. H., and Unser, M. (2017). A review of convolutional neural networks for inverse problems in imaging. *arXiv preprint arXiv:1710.04011*.
- Milletari, F., Navab, N., and Ahmadi, S.-A. (2016). V-net: Fully convolutional neural networks for volumetric medical image segmentation. In *3D Vision (3DV), 2016 Fourth International Conference on*, pages 565–571. IEEE.
- Morin, R., Basarab, A., and Kouamé, D. (2012). Alternating direction method of multipliers framework for super-resolution in ultrasound imaging. In *Biomedical Imaging (ISBI), 2012 9th IEEE International Symposium on*, pages 1595–1598. IEEE.
- Ng, M. K., Weiss, P., and Yuan, X. (2010). Solving constrained total-variation image restoration and reconstruction problems via alternating direction methods. *SIAM journal on Scientific Computing*, 32(5):2710–2736.
- Nguyen, N. and Milanfar, P. (2000). A wavelet-based interpolation-restoration method for super-resolution (wavelet superresolution). *Circuits, Systems and Signal Processing*, 19(4):321–338.
- Nishiyama, K. K. and Shane, E. (2013). Clinical imaging of bone microarchitecture with HR-pQCT. *Current osteoporosis reports*, 11(2):147–155.
- Nuzzo, S., Lafage-Proust, M., Martin-Badosa, E., Boivin, G., Thomas, T., Alexandre, C., and Peyrin, F. (2002). Synchrotron radiation microtomography allows the analysis of three-dimensional microarchitecture and degree of mineralization of human iliac crest biopsy specimens: Effects of etidronate treatment. *Journal of Bone and Mineral Research*, 17(8):1372–1382.
- Ochs, P., Chen, Y., Brox, T., and Pock, T. (2014). ipiano: Inertial proximal algorithm for nonconvex optimization. *SIAM Journal on Imaging Sciences*, 7(2):1388–1419.
- Odgaard, A. (1997). Three-dimensional methods for quantification of cancellous bone architecture. *Bone*, 20(4):315–328.

- Oktay, O., Ferrante, E., Kamnitsas, K., Heinrich, M., Bai, W., Caballero, J., Cook, S. A., de Marvao, A., Dawes, T., O'Regan, D. P., et al. (2018). Anatomically constrained neural networks (acnns): application to cardiac image enhancement and segmentation. *IEEE transactions on medical imaging*, 37(2):384–395.
- on Prevention, W. S. G., of Osteoporosis, M., and Organization, W. H. (2003). *Prevention and management of osteoporosis: report of a WHO scientific group*. Number 921. World Health Organization.
- Otsu, N. (1975). A threshold selection method from gray-level histograms. *Automatica*, 11(285-296):23–27.
- Parikh, N., Boyd, S. P., et al. (2014). Proximal Algorithms. *Foundations and Trends in optimization*, 1(3):127–239.
- Parzen, E. (1962). On estimation of a probability density function and mode. *The annals of mathematical statistics*, 33(3):1065–1076.
- Peyrin, F., Attali, D., Chappard, C., and Benhamou, C. L. (2010). Local plate/rod descriptors of 3D trabecular bone micro-CT images from medial axis topologic analysis. *Medical physics*, 37(8):4364–4376.
- Peyrin, F. and Engelke, K. (2012). CT imaging: Basics and new trends. In *Handbook of Particle Detection and Imaging*, pages 883–915. Springer.
- Peyrin, F., Mastrogiacomo, M., Cancedda, R., and Martinetti, R. (2007). SEM and 3D synchrotron radiation micro-tomography in the study of bioceramic scaffolds for tissue-engineering applications. *Biotechnology and bioengineering*, 97(3):638–648.
- Peyrin, F., Toma, A., Sixou, B., Denis, L., Burghardt, A., and Pialat, J.-B. (2015). Semi-blind joint super-resolution/segmentation of 3D trabecular bone images by a TV box approach. In *Signal Processing Conference (EUSIPCO), 2015 23rd European*, pages 2811–2815. IEEE.
- Pialat, J., Vilayphiou, N., Boutroy, S., Gouttenoire, P., Sornay-Rendu, E., Chapurlat, R., and Peyrin, F. (2012). Local topological analysis at the distal radius by HR-pQCT: application to in vivo bone microarchitecture and fracture assessment in the ofely study. *Bone*, 51(3):362–368.
- Polyak, B. T. (1964). Some methods of speeding up the convergence of iteration methods. *USSR Computational Mathematics and Mathematical Physics*, 4(5):1–17.
- Reddy, B. S. and Chatterji, B. N. (1996). An FFT-based technique for translation, rotation, and scale-invariant image registration. *IEEE transactions on image processing*, 5(8):1266–1271.
- Reeves, S. J. and Mersereau, R. M. (1992). Blur identification by the method of generalized cross-validation. *IEEE Transactions on Image Processing*, 1(3):301–311.
- Riggs, B. L. and Melton, L. J. (1983). *The American journal of medicine*, 75(6):899–901.
- Robinson, M. D., Toth, C. A., Lo, J. Y., and Farsiu, S. (2010). Efficient fourier-wavelet super-resolution. *IEEE Transactions on Image Processing*, 19(10):2669–2681.
- Rodan, G. A. (1992). Introduction to bone biology. *Bone*, 13:S3–S6.

- Ronneberger, O., Fischer, P., and Brox, T. (2015). U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, pages 234–241. Springer.
- Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1986). Learning representations by back-propagating errors. *nature*, 323(6088):533.
- Salomé, M., Peyrin, F., Cloetens, P., Odet, C., Laval-Jeantet, A.-M., Baruchel, J., and Spanne, P. (1999). A synchrotron radiation microtomography system for the analysis of trabecular bone samples. *Medical Physics*, 26(10):294–2204.
- Samson, C., Blanc-Féraud, L., Aubert, G., and Zerubia, J. (2000). A variational model for image classification and restoration. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22(5):460–472.
- Scherzer, O., Grasmair, M., Grossauer, H., Haltmeier, M., and Lenzen, F. (2009). *Variational methods in imaging*, volume 320. Springer.
- Schroder, H. M., Petersen, K. K., and Erlandsen, M. (1993). Occurrence and incidence of the second hip fracture. *Clinical orthopaedics and related research*, 289:166–169.
- Shamonin, D. P., Bron, E. E., Lelieveldt, B. P., Smits, M., Klein, S., and Staring, M. (2014). Fast parallel image registration on cpu and gpu for diagnostic classification of alzheimer’s disease. *Frontiers in neuroinformatics*, 7:50.
- Shi, H. and Ward, R. (2002). Canny edge based image expansion. In *Circuits and Systems, 2002. ISCAS 2002. IEEE International Symposium on*, volume 1, pages I–I. IEEE.
- Shi, W., Caballero, J., Huszár, F., Totz, J., Aitken, A. P., Bishop, R., Rueckert, D., and Wang, Z. (2016). Real-time single image and video super-resolution using an efficient sub-pixel convolutional neural network. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 1874–1883.
- Silman, A. J. (1995). The patient with fracture: the risk of subsequent fractures. *The American journal of medicine*, 98(2):12S–16S.
- Simonyan, K. and Zisserman, A. (2014). Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*.
- Soltani, S., Andersen, M. S., and Hansen, P. C. (2017). Tomographic image reconstruction using training images. *Journal of Computational and Applied Mathematics*, 313:243–258.
- Stone, H. S., Orchard, M. T., Chang, E.-C., and Martucci, S. A. (2001). A fast direct Fourier-based algorithm for subpixel registration of images. *IEEE Transactions on geoscience and remote sensing*, 39(10):2235–2243.
- Sun, L., Cho, S., Wang, J., and Hays, J. (2013). Edge-based blur kernel estimation using patch priors. In *Computational Photography (ICCP), 2013 IEEE International Conference on*, pages 1–8. IEEE.
- Szegedy, C., Ioffe, S., Vanhoucke, V., and Alemi, A. A. (2017). Inception-v4, inception-resnet and the impact of residual connections on learning. In *AAAI*, volume 4, page 12.

- Tai, Y., Yang, J., and Liu, X. (2017). Image super-resolution via deep recursive residual network. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, volume 1, page 5.
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, pages 267–288.
- Tjong, W., Kazakia, G. J., Burghardt, A. J., and Majumdar, S. (2012). The effect of voxel size on high-resolution peripheral computed tomography measurements of trabecular and cortical bone microstructure. *Medical physics*, 39(4):1893–1903.
- Tofighi, M., Yorulmaz, O., Köse, K., Yıldırım, D. C., Çetin-Atalay, R., and Çetin, A. E. (2016). Phase and tv based convex sets for blind deconvolution of microscopic images. *IEEE Journal of Selected Topics in Signal Processing*, 10(1):81–91.
- Toma, A., Denis, L., Sixou, B., Pialat, J.-B., and Peyrin, F. (2014a). Total variation super-resolution for 3D trabecular bone micro-structure segmentation. In *2014 22nd European Signal Processing Conference (EUSIPCO)*, pages 2220–2224. IEEE.
- Toma, A., Sixou, B., Denis, L., Pialat, J.-B., and Peyrin, F. (2014b). Higher order total variation super-resolution from a single trabecular bone image. In *Biomedical Imaging (ISBI), 2014 IEEE 11th International Symposium on*, pages 1152–1155. IEEE.
- Toma, A., Sixou, B., and Peyrin, F. (2015). Iterative choice of the optimal regularization parameter in TV image restoration. *Inverse Problems & Imaging*, 9(4).
- Tong, T., Li, G., Liu, X., and Gao, Q. (2017). Image super-resolution using dense skip connections. In *Computer Vision (ICCV), 2017 IEEE International Conference on*, pages 4809–4817. IEEE.
- Toriwaki, J. and Yonekura, T. (2002). Euler number and connectivity indexes of a three dimensional digital picture. *FORMA-TOKYO-*, 17(3):183–209.
- Tsai, R. (1984). Multiframe image restoration and registration. *Advance Computer Visual and Image Processing*, 1:317–339.
- Umehara, K., Ota, J., and Ishida, T. (2017). Application of super-resolution convolutional neural network for enhancing image resolution in chest CT. *Journal of digital imaging*, pages 1–10.
- Vese, L. A. and Chan, T. F. (2002). A multiphase level set framework for image segmentation using the Mumford and Shah model. *International journal of computer vision*, 50(3):271–293.
- Walker, M. D., McMahon, D. J., Udesky, J., Liu, G., and Bilezikian, J. P. (2009). Application of high-resolution skeletal imaging to measurements of volumetric bmd and skeletal microarchitecture in chinese-american and white women: explanation of a paradox. *Journal of Bone and Mineral Research*, 24(12):1953–1959.
- Wang, S., Ding, Z., and Fu, Y. (2018). Marginalized Denoising Dictionary Learning With Locality Constraint. *IEEE Transactions on Image Processing*, 27(1):500–510.
- Wang, S., Zhang, L., Liang, Y., and Pan, Q. (2012). Semi-coupled dictionary learning with applications to image super-resolution and photo-sketch synthesis. In *Computer Vision and Pattern Recognition (CVPR), 2012 IEEE Conference on*, pages 2216–2223. IEEE.

- Wang, Z. (1990). Interpolation using type i discrete cosine transform. *Electronics letters*, 26(15):1170–1172.
- Weaver, C., Gordon, C., Janz, K., Kalkwarf, H., Lappe, J., Lewis, R., O’Karma, M., Wallace, T., and Zemel, B. (2016). The National Osteoporosis Foundation’s position statement on peak bone mass development and lifestyle factors: a systematic review and implementation recommendations. *Osteoporosis International*, 27(4):1281–1386.
- Xu, L., Tao, X., and Jia, J. (2014). Inverse kernels for fast spatial deconvolution. In *European Conference on Computer Vision*, pages 33–48. Springer.
- Xu, Q., Yu, H., Mou, X., Zhang, L., Hsieh, J., and Wang, G. (2012). Low-dose X-ray CT reconstruction via dictionary learning. *IEEE Transactions on Medical Imaging*, 31(9):1682–1697.
- Yang, J., Wang, Z., Lin, Z., Cohen, S., and Huang, T. (2012). Coupled dictionary training for image super-resolution. *IEEE transactions on image processing*, 21(8):3467–3478.
- Yang, J., Wright, J., Huang, T. S., and Ma, Y. (2010a). Image super-resolution via sparse representation. *IEEE transactions on image processing*, 19(11):2861–2873.
- Yang, M., Zhang, L., Yang, J., and Zhang, D. (2010b). Metaface learning for sparse representation based face recognition. In *Image Processing (ICIP), 2010 17th IEEE International Conference on*, pages 1601–1604. IEEE.
- Yu, L., Yang, X., Chen, H., Qin, J., and Heng, P.-A. (2017). Volumetric convnets with mixed residual connections for automated prostate segmentation from 3D MR images. In *AAAI*, pages 66–72.
- Zeiler, M. D. and Fergus, R. (2014). Visualizing and understanding convolutional networks. In *European conference on computer vision*, pages 818–833. Springer.
- Zeyde, R., Elad, M., and Protter, M. (2010). On single image scale-up using sparse-representations. In *International conference on curves and surfaces*, pages 711–730. Springer.
- Zhang, K., Zuo, W., Chen, Y., Meng, D., and Zhang, L. (2017). Beyond a gaussian denoiser: Residual learning of deep cnn for image denoising. *IEEE Transactions on Image Processing*, 26(7):3142–3155.
- Zhang, Y., Tian, Y., Kong, Y., Zhong, B., and Fu, Y. (2018). Residual dense network for image super-resolution. In *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*.
- Zhao, C., Carass, A., Dewey, B. E., and Prince, J. L. (2018). Self super-resolution for magnetic resonance images using deep networks. In *Biomedical Imaging (ISBI 2018), 2018 IEEE 15th International Symposium on*, pages 365–368. IEEE.
- Zhao, X., Wu, Y., Tian, J., and Zhang, H. (2015). Single image super-resolution via blind blurring estimation and dictionary learning. In *CCF Chinese Conference on Computer Vision*, pages 22–33. Springer.
- Zou, H. and Hastie, T. (2005). Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 67(2):301–320.



FOLIO ADMINISTRATIF

THESE DE L'UNIVERSITE DE LYON OPEREE AU SEIN DE L'INSA LYON

NOM : LI
(avec précision du nom de jeune fille, le cas échéant)

DATE de SOUTENANCE : 20/12/2018

Prénoms : Yufei

TITRE : Joint super-resolution/segmentation approaches for the tomographic images analysis of the bone micro-structure

NATURE : Doctorat

Numéro d'ordre : 2018LYSEI125

Ecole doctorale : ELECTRONIQUE, ELECTROTECHNIQUE, AUTOMATIQUE

Spécialité : Traitement du signal et de l'image

RESUME : L'ostéoporose est une maladie caractérisée par une perte de masse osseuse et la dégradation de la micro-architecture osseuse. La micro-architecture joue un rôle important pour le diagnostic de l'ostéoporose. La tomographie périphérique à haute résolution (HR-pQCT) est actuellement l'une des meilleures techniques de tomodensitométrie pour l'étude de la micro-architecture osseuse en 3D. Elle permet d'accéder à la mesure de la micro-architecture osseuse, mais la résolution est limitée. La Micro-CT, qui peut fournir des images à une résolution spatiale beaucoup plus élevée, se limite à un examen ex vivo. Dans cette thèse, nous tentons de résoudre un problème de super résolution/segmentation joint, afin d'améliorer la qualité des images HR-pQCT. Les images HR-pQCT seront considérées comme des images à basse résolution et les images micro-CT du même spécimen avec une taille de voxel de 24 μ m, seront considérées comme des images à haute résolution et utilisées comme une vérité-terrain pour la micro-architecture osseuse. Les deux premiers chapitres sont consacrés à la présentation du contexte médical et des fondements mathématiques. Au chapitre 3, nous présentons notre première contribution : la détermination du noyau de dégradation. Dans le chapitre 4, nous avons utilisé des approches variationnelles basées sur la régularisation de la variation totale et le potentiel de double puits. Comme la structure de l'os trabéculaire n'étant pas bien récupérée, nous proposons d'utiliser la méthode d'apprentissage par dictionnaire pour préserver les détails structuraux. Dans les chapitres 6 et 7, nous avons utilisé des approches d'apprentissage approfondi pour affiner notre travail présenté au chapitre 5. Pendant la thèse, nous avons essayé différentes approches et la qualité de l'image s'améliore progressivement. Notre résultats montre que la méthode d'apprentissage profond est très prometteur pour l'application clinique dans le futur.

MOTS-CLÉS : Bone micro-architecture, CT images, super resolution

Laboratoire (s) de recherche : CREATIS

Directeur de thèse: Bruno SIXOU

Président de jury : TALBOT Hugues, professeur associé, ESIEE

Composition du jury :

KOUEME Denis,
SOUSSEN Charles,
CHAPPARD Christine,
TALBOT Hugues,
SIXOU Bruno,
PEYRIN Françoise,

Professeur des Universités,
Professeur des Universités,
CR1, Inserm,
professeur associé,
Maitre de conférence, HDR
DR Inserm,

IRIT,
Centrale-Supelec
B2OA
ESIEE
INSA-LYON
INSA-LYON, ESRF

Rapporteur
Rapporteur
Examinatrice
Examineur
Directeur de thèse
Co-directrice de thèse